



Nagios® Version 2.x  
Documentation

Copyright © 1999-2006 Ethan Galstad  
[www.nagios.org](http://www.nagios.org)

Last Updated: 11-27-2006

[ [Table of Contents](#) ]

Nagios and the Nagios logo are registered trademarks of Ethan Galstad. All other trademarks, servicemarks, registered trademarks, and registered servicemarks mentioned herein may be the property of their respective owner(s). The information contained herein is provided AS IS with NO WARRANTY OF ANY KIND, INCLUDING THE WARRANTY OF DESIGN, MERCHANTABILITY, AND FITNESS FOR A PARTICULAR PURPOSE.



Version 2.0 Documentation

Table of Contents

---

## About

- [What is Nagios?](#)
- [System requirements](#)
- [Licensing](#)
- [Downloading the latest version](#)
- [Other monitoring utilities](#)

## Release Notes

- [What's new in this version](#)
- [Change log](#)

## Support

- [Self-service and commercial support](#)

## Getting Started

- [Advice for beginners](#)

## Installing Nagios

- [Compiling and installing Nagios](#)
- [Setting up the web interface](#)

## Configuring Nagios

- [Configuration overview](#)
- [Main configuration file options](#)
- [Object configuration file options](#)
- [CGI configuration file options](#)
- [Configuring authorization for the CGIs](#)

## Running Nagios

- [Verifying the configuration](#)
- [Starting Nagios](#)
- [Stopping and restarting Nagios](#)

## Nagios Plugins

- [Standard plugins](#)
- [Writing your own plugins](#)

## Nagios Addons

[NRPE](#) - Daemon and plugin for executing plugins on remote hosts

[NSCA](#) - Daemon and client program for sending passive check results across the network

## Theory Of Operation

[Determining status and reachability of network hosts](#)

[Network outages](#)

[Notifications](#)

[Plugin theory](#)

[Service check scheduling](#)

[State types](#)

[Time periods](#)

## Advanced Topics

[Event handlers](#)

[External commands](#)

[Indirect host and service checks](#)

[Passive service checks](#)

[Volatile services](#)

[Service and host result freshness checks](#)

[Distributed monitoring](#)

[Redundant and failover monitoring](#)

[Detection and handling of state flapping](#)

[Service check parallelization](#)

[Notification escalations](#)

[Monitoring service and host clusters](#)

[Host and service dependencies](#)

[State stalking](#)

[Performance data](#)

[Scheduled host and service downtime](#)

[Using the embedded Perl interpreter](#)

[Adaptive monitoring](#)

[Object inheritance](#)

[Time-saving tips for object definitions](#)

## Integration With Other Software

[SNMP Traps](#)

[TCP Wrappers](#)

## Miscellaneous

[Securing Nagios](#)

[Tuning Nagios for maximum performance](#)

[Using the nagiosstats utility](#)

[Using macros in commands](#)

[Information on the CGIs](#)

[Custom CGI headers and footers](#)

# About Nagios®

---

## What Is This?

Nagios® is a system and network monitoring application. It watches hosts and services that you specify, alerting you when things go bad and when they get better.

Nagios was originally designed to run under [Linux](#), although it should work under most other unices as well.

Some of the many features of Nagios® include:

- Monitoring of network services (SMTP, POP3, HTTP, NNTP, PING, etc.)
- Monitoring of host resources (processor load, disk usage, etc.)
- Simple plugin design that allows users to easily develop their own service checks
- Parallelized service checks
- Ability to define network host hierarchy using "parent" hosts, allowing detection of and distinction between hosts that are down and those that are unreachable
- Contact notifications when service or host problems occur and get resolved (via email, pager, or user-defined method)
- Ability to define event handlers to be run during service or host events for proactive problem resolution
- Automatic log file rotation
- Support for implementing redundant monitoring hosts
- Optional web interface for viewing current network status, notification and problem history, log file, etc.

## System Requirements

*The only requirement of running Nagios is a machine running Linux (or UNIX variant) and a C compiler. You will probably also want to have TCP/IP configured, as most service checks will be performed over the network.*

You are *not required* to use the CGIs included with Nagios. However, if you do decide to use them, you will need to have the following software installed...

1. A web server (preferably [Apache](#))
2. Thomas Boutell's [gd library](#) version 1.6.3 or higher (required by the [statusmap](#) and [trends](#) CGIs)

## Licensing

Nagios® is licensed under the terms of the [GNU General Public License](#) Version 2 as published by the [Free Software Foundation](#). This gives you legal permission to copy, distribute and/or modify Nagios under certain conditions. Read the 'LICENSE' file in the Nagios distribution or read the [online version of the license](#) for more details.

Nagios® is provided AS IS with NO WARRANTY OF ANY KIND, INCLUDING THE WARRANTY OF DESIGN, MERCHANTABILITY, AND FITNESS FOR A PARTICULAR PURPOSE.

## Acknowledgements

Several people have contributed to Nagios by either reporting bugs, suggesting improvements, writing plugins, etc. A list of some of the many contributors to the development of Nagios can be found at <http://www.nagios.org>.

### **Downloading The Latest Version**

You can check for new versions of Nagios at <http://www.nagios.org>.

---

Nagios and the Nagios logo are trademarks of Ethan Galstad. All other trademarks, servicemarks, registered trademarks, and registered servicemarks may be the property of their respective owner(s).

# What's New in Version 2.0

---

**Important:** Make sure you read through the documentation before sending a question to the mailing lists.

## Change Log

The change log for Nagios can be found online at <http://www.nagios.org/development/changelog.php> or in the **Changelog** file in the root directory of the source code distribution.

## Known Issues

There is a known issue that can affect Nagios 2.0 on FreeBSD systems. Hopefully this problem can be fixed in a 2.x release...

1. **FreeBSD and threads.** On FreeBSD there's a native user-level implementation of threads called 'pthread' and there's also an optional ports collection 'linuxthreads' that uses kernel hooks. Some folks from Yahoo! have reported that using the pthread library causes Nagios to pause under heavy I/O load, causing some service check results to be lost. Switching to linuxthreads seems to help this problem, but not fix it. The lock happens in liblthread's `__pthread_acquire()` - it can't ever acquire the spinlock. It happens when the main thread forks to execute an active check. On the second fork to create the grandchild, the grandchild is created by fork, but never returns from liblthread's fork wrapper, because it's stuck in `__pthread_acquire()`. Maybe some FreeBSD users can help out with this problem.

## Changes and New Features

1. **Macro Changes** - [Macros](#) have undergone a major overhaul. You will have to update most of your command definitions to match the new macros. Most macros are now available as environment variables. Also, "on-demand" host and service macros have been added. See the [documentation on macros](#) for more information.
2. **Hostgroup Changes**
  - **Hostgroup escalations removed** - Hostgroup escalations have been removed. Their functionality can be duplicated by using the `hostgroup_name` directive in [hostgroup definitions](#).
  - **Member directive changes** - Hostgroup [definitions](#) can now contain multiple *members* directives, which should make editing the config files easier when you have a lot of member hosts. Alternatively, you may use the `hostgroups` directive in [host definitions](#) to specify what hostgroup(s) a particular host is a member of.
  - **Contact group changes** - The `contact_groups` directive has been moved from hostgroup definitions to [host definitions](#). This was done in order to maintain consistency with the way service contacts are specified. Make sure to update your config files!
  - **Authorization changes** - Authorization for access to hostgroups in the CGIs has been changed. You must now be authorized for all hosts that are members of the hostgroup in order to be authorized for the hostgroup.
3. **Host Changes**
  - **Host freshness checking** - Freshness checking has been added for host checks. This is controlled by the [check\\_host\\_freshness](#) option, along with the `check_freshness` directive in [host definitions](#).
  - **OCHP Command** - Host checks can now be obsessed over, just as services can be. The [OCHP command](#) is run for all hosts that have the `obsess_over_host` directive enabled in their [host definition](#).
4. **Host Check Changes**
  - **Regularly scheduled checks** - You can now schedule regular checks of hosts by using the

*check\_interval* directive in [host definitions](#). **NOTE:** Listen up! You should use regularly scheduled host checks rather sparingly. They are not necessary for normal operation (on-demand checks are already performed when necessary) and can [negatively affect performance](#) if used improperly. You've been warned.

- **Passive host checks** - Passive host checks are now supported if you've enabled them with the [accept\\_passive\\_host\\_checks](#) option in the main config file and the *accept\_passive\_host\_checks* directive in the [host definition](#). [Passive host checks](#) can make setting up [redundant](#) or [distributed](#) monitoring environments easier. **NOTE:** There are some problems with passive host checks that you should be aware of - read more about them [here](#).

## 5. Retention Changes

- **Retention of scheduling information** - Host and service check scheduling information (next check times) can now be retained across program restarts using the [use\\_retained\\_scheduling\\_info](#) directive.
- **Smarter retention** - Values of various host and service directives that can be retained across program restarts are now only retained if they are changed during runtime by an [external command](#). This should make things less confusing to people when they try and modify host and service directive values and then restart Nagios, expecting to see some changes.
- **More stuff retained** - More information is now retained across program restarts, including [flap detection](#) history. Hoorah!

## 6. Extended Info Changes

- **New location** - Extended [host info](#) and [service info](#) definitions are now stored in object config files along with host definitions, etc. As a result, extended info definitions are now parsed and validated by the Nagios daemon before startup.
- **New directives** - Extended [host info](#) and [service info](#) definitions now have two new directives: *notes* and *action\_url*.

## 7. Embedded Perl Changes

- **p1.pl location** - You can now specify the location of the embedded Perl "helper" file (p1.pl) using the [p1\\_file](#) directive.

## 8. Notification Changes

- **Flapping notifications** - Notifications are now sent out when [flapping](#) starts and stops for hosts and services. This feature can be controlled using the *f* option in the *notification\_options* for [contacts](#), [hosts](#) and [services](#).
- **Better logic** - Notification logic has been improved a bit. This should prevent recovery notifications getting sent out when no problem notification was sent out to begin with.
- **Service notifications** - Before service notifications are sent out, notification [dependencies](#) for the host are now checked. If host notifications are not deemed to be viable, notifications for the service will not be sent out either.
- **Escalation options** - Time period and state options have been added to [host](#) and [service](#) escalations. This gives you more control in determining when escalations can be used. More information on escalations can be found [here](#).

- 9. **Service Groups Added** - [Service groups](#) have now been added. They allow you to group services together for display purposes in the CGIs and [can be referenced](#) in service dependency and service escalation definitions to make configuration a bit easier.

- 10. **Triggered Downtime Added** - Support for what's called "triggered" downtime has been added for host and service downtime. Triggered downtime allows you to define downtime that should start at the same time another downtime starts (very useful for scheduling downtime for child hosts when the parent host is scheduled for flexible downtime). More information on triggered downtime can be found [here](#).

- 11. **New Stats Utility** - A new utility called 'nagiosstats' is now included in the Nagios distribution. It's a command-line utility that allows you to view current statistics for a running Nagios process. It can also produce data compatible with MRTG, so you can graph statistical information. More information on how to use the utility can be found [here](#).

12. **Adaptive Monitoring** - You can now change certain attributes relating to host and service checks (check command, check interval, max check attempts etc.) during runtime by submitting the appropriate external commands. This kind of adaptive monitoring will probably not be of much use to the majority of users out there, but it does provide a way for doing some neat stuff. More information on adaptive monitoring can be found [here](#).
13. **Performance Data Changes** - The methods for processing performance data have changed slightly. You can now process performance data by executing external commands and/or writing to files without recompiling Nagios. Read the documentation on [performance data](#) for more information.
14. **Native DB Support Dropped** - Native support for storing various types of data (status, retention, comment, downtime, etc.) in MySQL and PostgreSQL has been dropped. Stop whining. I expect someone will develop an alternative using the new event broker sometime in the near future. Besides, DB support was not well implemented and dropping native DB support will make things easier for newbies to understand (one less thing to figure out).
15. **Event Broker API** - An API has been created to allow individual developers to create addons that integrate with the core Nagios daemon. Documentation on the event broker API will be created as the 2.x code matures and will be available on the Nagios website.
16. **Misc Changes**
  - **All commands can contain arguments** - All command types (host checks, notifications, performance data processors, event handlers, etc.) can contain arguments (seperated from the command name by ! characters). Arguments are substituted in the command line using [\\$ARGx macros](#).
  - **Config directory recursion** - Nagios now recursively processes all config files found in subdirectories of the directories specified by the [cfg\\_dir directive](#).
  - **Old config file support dropped** - Support for older (non-template) style object and extended info config files has been dropped.
  - **Faster searches** - Objects are now stored in a chained hash in order to speed searches. This should greatly improve the performance of the CGIs.
  - **Worker threads** - A few worker threads have been added in order to artificially buffer data for the [external command file](#) and the internal pipe used to process service check results. This should substantially increase performance in larger setups.
  - **Logging changes** - Initial host and service states are now logged a bit differently. Also, the initial states of all hosts and services are logged immediately after all [log rotations](#). This should help with all those "undetermined time" problems in the availability and trends CGIs.
  - **Cached object config file** - An [object cache file](#) is now created by Nagios at startup. It should help speed up the CGIs a bit and allow you to edit you object config files while Nagios is running without affecting the CGI output.
  - **Initial check limits** - You can now specify timeframes in which the initial checks of all hosts and services should be performed after Nagios start. These timeframes are controlled by the [max\\_host\\_check\\_spread](#) and [max\\_service\\_check\\_spread](#) variables.
  - **"Sticky" acknowledgements** - You can now designate host and service acknowledgements as being "sticky" or not. Sticky acknowledgements suppress notifications until a host or service fully recovers to an UP or OK state. Non-sticky acknowledgements only suppress notifications until a host or service changes state.
  - **Changed in checking clusters** - The way you monitor service and host "clusters" has now changed and is more reliable than before. This is due to the incorporation of on-demand macros and a new plugin (check\_cluster2). Read more about checking clusters [here](#).
  - **Regular expression matching** - Regular expression matching of various object directives can be enabled using the [use\\_regexp\\_matching](#) and [use\\_true\\_regexp\\_matching](#) variables. Information on how and where regular expression matching can be used can be found in the [template tips and tricks](#) documentation.
  - **Service pseudo-states** - Support for some redundant service pseudo-states have been removed from the status CGI. This will affect any hardcoded URLs which use the servicestatustypes=X parameter for the CGI. Check include/statusdata.h for the new list of service states that you can



use.

- **Freshness check changes** - The freshness check logic has been changed slightly. Freshness checks will not occur if the current time is not valid for the host or service *check\_timeperiod*. Also, freshness checks will no longer occur if both the host or service *check\_interval* and *freshness\_threshold* variables are set to zero (0).
-

## Advice for Beginners

---

Congrats on choosing to try Nagios! Nagios is quite powerful and flexible, but unfortunately its not very friendly to newbies. Why? Because it takes a lot of work to get it installed and configured properly. That being said, if you stick with it and manage to get it up and running, you'll never want to be without it. :-) Here are some very important things to keep in mind for those of you who are first-time users of Nagios:

1. **Relax - its going to take some time.** Don't expect to be able to compile Nagios and start it up right off the bat. Its not that easy. In fact, its pretty difficult. If you don't want to spend time learning how things work and getting things running smoothly, don't bother using this software. Instead, pay someone to monitor your network for you or hire someone to install Nagios for you. :-)
  2. **Read the documentation.** Nagios is difficult enough to configure when you've got a good grasp of what's going on, and nearly impossible if you don't. Do yourself a favor and read before blindly attempting to install and run Nagios. If you're the type who doesn't want to take the time to read the documentation, you'll probably find that others won't find the time to help you out when you have problems. RTFM.
  3. **Use the sample config files.** Sample configuration files are provided with Nagios. Look at them, modify them for your particular setup and test them! The sample files are just that - samples. There's a very good chance that they won't work for you without modifications. Sample config files can be found in the *sample-config/* subdirectory of the Nagios distribution.
  4. **Seek the help of others.** If you've read the documentation, reviewed the sample config files, and are still having problems, try sending a *descriptive* email message describing your problems to the *nagios-users* mailing list. Due to the amount of work that I have to do for this project, I am unable to answer most of the questions that get sent directly to me, so your best source of help is going to be the mailing list. If you've done some background reading and you provide a good problem description, odds are that someone will give you some pointers on getting things working properly.
-

# Installing Nagios

---

**Important:** Installing and configuring Nagios is rather involved. You can't just compile the binaries, run the program and sit back. There's a lot to setup before you can actually start monitoring anything. Relax, take your time and read all the documentation - you're going to need it. Okay, let's get started...

## Become Root

You'll need to have root access to install Nagios as described in this documentation, as you'll be creating users and group, modifying your web server config files, etc. Either login as root before you begin or use the `su` command to change to the root user from another account.

## Getting The Latest Version

You can download the latest version of Nagios from <http://www.nagios.org/download>.

## Unpacking The Distribution

To unpack the Nagios distribution, use the following command:

```
tar xzf nagios-version.tar.gz
```

When you have finished executing these commands, you should find a **nagios-version** directory that has been created in your current directory. Inside that directory you will find all the files that comprise the core Nagios distribution.

## Create Nagios User/Group

You're probably going to want to run Nagios under a normal user account, so add a new user (and group) to your system with the following command (this will vary depending on what OS you're running):

```
adduser nagios
```

## Create Installation Directory

Create the base directory where you would like to install Nagios as follows...

```
mkdir /usr/local/nagios
```

Change the owner of the base installation directory to be the Nagios user and group you added earlier as follows:

```
chown nagios.nagios /usr/local/nagios
```

## Identify Web Server User

You're probably going to want to issue **external commands** (like acknowledgements and scheduled downtime) from the web interface. To do so, you need to identify the user your web server runs as (typically *apache*, although this may differ on your system). This setting is found in your web server configuration file. The following command can be used to quickly determine what user Apache is running as (paths may differ on your system):

```
grep "^User" /etc/httpd/conf/httpd.conf
```

### Add Command File Group

Next we're going to create a new group whose members include the user your web server is running as and the user Nagios is running as. Let's say we call this new group '**nagcmd**' (you can name it differently if you wish). On RedHat Linux you can use the following command to add a new group (other systems may differ):

```
/usr/sbin/groupadd nagcmd
```

Next, add the users that your web server and Nagios run as to the newly created group with the following commands (I'll assume *apache* and *nagios* are the respective users):

```
/usr/sbin/usermod -G nagcmd apache  
/usr/sbin/usermod -G nagcmd nagios
```

### Run the Configure Script

Run the configure script to initialize variables and create a Makefile as follows...(the last two options: `--with-command-xxx` are optional, but needed if you want to issue [external commands](#))

```
./configure --prefix=prefix --with-cgiurl=cgiurl --with-htmurl=htmurl --with-nagios-user=someuser  
--with-nagios-group=somegroup --with-command-group=cmdgroup
```

- Replace *prefix* with the installation directory that you created in the step above (default is */usr/local/nagios*)
- Replace *cgiurl* with the actual url you will be using to access the [CGIs](#) (default is */nagios/cgi-bin*). Do NOT append a slash at the end of the url.
- Replace *htmurl* with the actual url you will be using to access the HTML for the main interface and documentation (default is */nagios/*)
- Replace *someuser* with the name of a user on your system that will be used for setting permissions on the installed files (default is *nagios*)
- Replace *somegroup* with the name of a group on your system that will be used for setting permissions on the installed files (default is *nagios*)
- Replace *cmdgroup* with the name of the group running the web server (default is *nagios*, in the example above it was *nagcmd*). This will allow group members (i.e. your web server) to be able to submit external commands to Nagios.

### Compile Binaries

Compile Nagios and the CGIs with the following command:

```
make all
```

### Installing The Binaries And HTML Files

Install the binaries and HTML files (documentation and main web page) with the following command:

```
make install
```

### Installing An Init Script

If you wish, you can also install the sample init script to */etc/rc.d/init.d/nagios* with the following command:

## make install-init

You may have to edit the init script to make sense with your particular OS and Nagios installation by editing paths, etc.

### Directory Structure And File Locations

Change to the root of your Nagios installation directory with the following command...

**cd /usr/local/nagios**

You should see five different subdirectories. A brief description of what each directory contains is given in the table below.

| Sub-Directory | Contents                                                                                                                                  |
|---------------|-------------------------------------------------------------------------------------------------------------------------------------------|
| bin/          | Nagios core program                                                                                                                       |
| etc/          | <a href="#">Main</a> , <a href="#">resource</a> , <a href="#">object</a> , and <a href="#">CGI</a> configuration files should be put here |
| sbin/         | <a href="#">CGIs</a>                                                                                                                      |
| share/        | HTML files (for web interface and online documentation)                                                                                   |
| var/          | Empty directory for the <a href="#">log file</a> , <a href="#">status file</a> , <a href="#">retention file</a> , etc.                    |
| var/archives  | Empty directory for the <a href="#">archived logs</a>                                                                                     |
| var/rw        | Empty directory for the <a href="#">external command file</a>                                                                             |

### Installing The Plugins

In order for Nagios to be of any use to you, you're going to have to download and install some [plugins](#). Plugins are usually installed in the **libexec/** directory of your Nagios installation (i.e. */usr/local/nagios/libexec*). Plugins are scripts or binaries which perform all the service and host checks that constitute monitoring. You can grab the latest release of the plugins from the [Nagios downloads page](#) or directly from the [SourceForge project page](#).

### Setup The Web Interface

You're probably going to want to use the web interface, so you'll also have to read the instructions on [setting up the web interface](#) and configuring web authentication, etc. next.

### Configuring Nagios

So now you have things compiled and installed, but you still haven't configured Nagios or defined objects (hosts, services, etc.) that should be monitored. Information on configuring Nagios and defining objects can be found [here](#). There's a lot to configure, but don't let it discourage you - its well worth it.

---

# Setting Up The Web Interface

---

## Notes

In these instructions I will assume that you are running the [Apache](#) web server on your machine. If you are using some other web server, you'll have to make changes where appropriate. I am also assuming that you used the `/usr/local/nagios` as the installation prefix.

## Sample Configuration

A sample Apache config file snippet is created when you run the configure script - you can find the sample config file (named `httpd.conf`) in the `sample-config/` subdirectory of the Nagios distribution. You will need to add the contents of this file to your Apache configuration files before you can access the Nagios web interface. The instructions found below detail how to manually add the appropriate configuration entries to Apache.

## Configure Aliases and Directory Options For The Web Interface

First you'll need to create appropriate entries for the Nagios web interface (HTML and CGIs) in your web server config file. Add the following snippet to your web server configuration file (i.e. `httpd.conf`), changing it to match any directory differences on your system.

```
ScriptAlias /nagios/cgi-bin /usr/local/nagios/sbin

<Directory "/usr/local/nagios/sbin">
    Options ExecCGI
    AllowOverride None
    Order allow,deny
    Allow from all
    AuthName "Nagios Access"
    AuthType Basic
    AuthUserFile /usr/local/nagios/etc/htpasswd.users
    Require valid-user
</Directory>

Alias /nagios /usr/local/nagios/share

<Directory "/usr/local/nagios/share">
    Options None
    AllowOverride None
    Order allow,deny
    Allow from all
    AuthName "Nagios Access"
    AuthType Basic
    AuthUserFile /usr/local/nagios/etc/htpasswd.users
    Require valid-user
</Directory>
```

**Note:** The default Nagios installation expects to find the HTML files and CGIs at `http://yourmachine/nagios/` and `http://yourmachine/nagios/cgi-bin/`, respectively. These locations can be changed using the `--with-htmurl` and `--with-cgiurl` options in the configure script.

**Important!** If you are installing Nagios on a multi-user system, you may want use [CGIWrap](#) to provide additional security between the CGIs and the [external command file](#). If you decide to use CGIWrap, the ScriptAlias you'll end up using will most likely be different from that mentioned above. More information on doing this can be found [here](#).

## Restart The Web Server

Once you've finished editing the Apache configuration file, you'll need to restart the web server with a command like this...

**/etc/rc.d/init.d/httpd restart**

## Configure Web Authentication

Once you have installed the web interface properly, you'll need to specify who can access the Nagios web interface. Follow [these instructions](#) to do this.

## Verify Your Changes

Don't forget to check and see if the changes you made to Apache work. You should be able to point your web browser at **http://yourmachine/nagios/** and get the web interface for Nagios. The CGIs may not display any information, but this will be remedied once you configure everything and start Nagios.

---

# Configuring Nagios

---

## Configuration Overview

There are several different configuration files that you're going to need to create or edit before you start monitoring anything. They are described below...

## Main Configuration File

The main configuration file (usually `/usr/local/nagios/etc/nagios.cfg`) contains a number of directives that affect how Nagios operates. This config file is read by both the Nagios process and the CGIs. This is the first configuration file you're going to want to create or edit.

Documentation for the main configuration file can be found [here](#).

A sample main configuration file is generated automatically when you run the **configure** script before compiling the binaries. Look for it either in the distribution directory or the `etc/` subdirectory of your installation. When you [install](#) the sample config files using the **make install-config** command, a sample main configuration file will be placed into your settings directory (usually `/usr/local/nagios/etc`). The default name of the main configuration file is **nagios.cfg**.

## Resource File(s)

Resource files can be used to store user-defined [macros](#). Resource files can also contain other information (like database connection settings), although this will depend on how you've compiled Nagios. The main point of having resource files is to use them to store sensitive configuration information and not make them available to the CGIs.

You can specify one or more optional resource files by using the `resource_file` directive in the [main configuration file](#).

## Object Definition Files

Object definition files are used to define hosts, services, hostgroups, contacts, contactgroups, commands, etc. This is where you define what things you want monitor and how you want to monitor them.

Documentation for the object definition files can be found [here](#).

## CGI Configuration File

The CGI configuration file (usually `/usr/local/nagios/etc/cgi.cfg`) contains a number of directives that affect the operation of the [CGIs](#).

Documentation for the CGI configuration file can be found [here](#).

A sample CGI configuration file is generated automatically when you run the **configure** script before compiling the binaries. When you [install](#) the sample config files using the **make install-config** command, the CGI configuration file will be placed in the same directory as the main and host config files (usually `/usr/local/nagios/etc`). The default name of the CGI configuration file is **cgi.cfg**.

---



# Main Configuration File Options

---

## Notes

When creating and/or editing configuration files, keep the following in mind:

1. Lines that start with a '#' character are taken to be comments and are not processed
2. Variables names must begin at the start of the line - no white space is allowed before the name
3. Variable names are case-sensitive

## Sample Configuration

A sample main configuration file is created in the base directory of the Nagios distribution when you run the configure script. The default name of the main configuration file is **nagios.cfg** - its usually placed in the **etc/** subdirectory of you Nagios installation (i.e. */usr/local/nagios/etc/*).

## Index

[Log file](#)  
[Object configuration file](#)  
[Object configuration directory](#)  
[Object cache file](#)  
[Resource file](#)  
[Temp file](#)  
[Status file](#)  
[Aggregated status updates option](#)  
[Aggregated status data update interval](#)  
[Nagios user](#)  
[Nagios group](#)  
[Notifications option](#)  
[Service check execution option](#)  
[Passive service check acceptance option](#)  
[Host check execution option](#)  
[Passive host check acceptance option](#)  
[Event handler option](#)  
[Log rotation method](#)  
[Log archive path](#)  
[External command check option](#)  
[External command check interval](#)  
[External command file](#)  
[Comment file](#)  
[Downtime file](#)  
[Lock file](#)  
[State retention option](#)  
[State retention file](#)  
[Automatic state retention update interval](#)  
[Use retained program state option](#)  
[Use retained scheduling info option](#)  
[Syslog logging option](#)  
[Notification logging option](#)  
[Service check retry logging option](#)  
[Host retry logging option](#)  
[Event handler logging option](#)

Initial state logging option  
External command logging option  
Passive check logging option  
Global host event handler  
Global service event handler  
Inter-check sleep time  
Service inter-check delay method  
Maximum service check spread  
Service interleave factor  
Maximum concurrent service checks  
Service reaper frequency  
Host inter-check delay method  
Maximum host check spread  
Timing interval length  
Auto-rescheduling option  
Auto-rescheduling interval  
Auto-rescheduling window  
Aggressive host checking option  
Flap detection option  
Low service flap threshold  
High service flap threshold  
Low host flap threshold  
High host flap threshold  
Soft service dependencies option  
Service check timeout  
Host check timeout  
Event handler timeout  
Notification timeout  
Obsessive compulsive service processor timeout  
Obsessive compulsive host processor timeout  
Performance data processor command timeout  
Obsess over services option  
Obsessive compulsive service processor command  
Obsess over hosts option  
Obsessive compulsive host processor command  
Performance data processing option  
Host performance data processing command  
Service performance data processing command  
Host performance data file  
Service performance data file  
Host performance data file template  
Service performance data file template  
Host performance data file mode  
Service performance data file mode  
Host performance data file processing interval  
Service performance data file processing interval  
Host performance data file processing command  
Service performance data file processing command  
Orphaned service check option  
Service freshness checking option  
Service freshness check interval  
Host freshness checking option  
Host freshness check interval  
Date format

[Illegal object name characters](#)  
[Illegal macro output characters](#)  
[Regular expression matching option](#)  
[True regular expression matching option](#)  
[Administrator email address](#)  
[Administrator pager](#)

## Log File

Format: **log\_file=<file\_name>**

Example: **log\_file=/usr/local/nagios/var/nagios.log**

This variable specifies where Nagios should create its main log file. This should be the first variable that you define in your configuration file, as Nagios will try to write errors that it finds in the rest of your configuration data to this file. If you have [log rotation](#) enabled, this file will automatically be rotated every hour, day, week, or month.

## Object Configuration File

Format: **cfg\_file=<file\_name>**

Example: **cfg\_file=/usr/local/nagios/etc/hosts.cfg**  
**cfg\_file=/usr/local/nagios/etc/services.cfg**  
**cfg\_file=/usr/local/nagios/etc/commands.cfg**

This directive is used to specify an [object configuration file](#) containing object definitions that Nagios should use for monitoring. Object configuration files contain definitions for hosts, host groups, contacts, contact groups, services, commands, etc. You can separate your configuration information into several files and specify multiple **cfg\_file=** statements to have each of them processed.

## Object Configuration Directory

Format: **cfg\_dir=<directory\_name>**

Example: **cfg\_dir=/usr/local/nagios/etc/commands**  
**cfg\_dir=/usr/local/nagios/etc/services**  
**cfg\_dir=/usr/local/nagios/etc/hosts**

This directive is used to specify a directory which contains [object configuration files](#) that Nagios should use for monitoring. All files in the directory with a **.cfg** extension are processed as object config files. Additionally, Nagios will recursively process all config files in subdirectories of the directory you specify here. You can separate your configuration files into different directories and specify multiple **cfg\_dir=** statements to have all config files in each directory processed.

## Object Cache File

Format: **object\_cache\_file=<file\_name>**

Example: **object\_cache\_file=/usr/local/nagios/var/objects.cache**

This directive is used to specify a file in which a cached copy of [object definitions](#) should be stored. The cache file is (re)created every time Nagios is (re)started and is used by the CGIs. It is intended to speed up config file caching in the CGIs and allow you to edit the source [object config files](#) while Nagios is running without affecting the output displayed in the CGIs.

### Resource File

Format: **resource\_file=<file\_name>**

Example: **resource\_file=/usr/local/nagios/etc/resource.cfg**

This is used to specify an optional resource file that can contain \$USERn\$ [macro](#) definitions. \$USERn\$ macros are useful for storing usernames, passwords, and items commonly used in command definitions (like directory paths). The CGIs will *not* attempt to read resource files, so you can set restrictive permissions (600 or 660) on them to protect sensitive information. You can include multiple resource files by adding multiple resource\_file statements to the main config file - Nagios will process them all. See the sample resource.cfg file in the base of the Nagios directory for an example of how to define \$USERn\$ macros.

### Temp File

Format: **temp\_file=<file\_name>**

Example: **temp\_file=/usr/local/nagios/var/nagios.tmp**

This is a temporary file that Nagios periodically creates to use when updating comment data, status data, etc. The file is deleted when it is no longer needed.

### Status File

Format: **status\_file=<file\_name>**

Example: **status\_file=/usr/local/nagios/var/status.dat**

This is the file that Nagios uses to store the current status of all monitored services. The status of all hosts associated with the service you monitor are also recorded here. This file is used by the CGIs so that current monitoring status can be reported via a web interface. The CGIs must have read access to this file in order to function properly. This file is deleted every time Nagios stops and recreated when it starts.

### Aggregated Status Updates Option

Format: **aggregate\_status\_updates=<0/1>**

Example: **aggregate\_status\_updates=1**

This option determines whether or not Nagios will aggregate updates of host, service, and program status data. If you do not enable this option, status data is updated every time a host or service checks occurs. This can result in high CPU loads and file I/O if you are monitoring a lot of services. If you want Nagios to only update status data (in the [status file](#)) every few seconds (as determined by the [status\\_update\\_interval](#) option), enable this option. If you want immediate updates, disable it. I would

highly recommend using aggregated updates (even at short intervals) unless you have good reason not to. Values are as follows:

- 0 = Disable aggregated updates
- 1 = Enabled aggregated updates (default)

### Aggregated Status Update Interval

Format: **status\_update\_interval=<seconds>**

Example: **status\_update\_interval=15**

This setting determines how often (in seconds) that Nagios will update status data in the [status file](#). The minimum update interval is five seconds. If you have disabled aggregated status updates (with the [aggregate\\_status\\_updates](#) option), this option has no effect.

### Nagios User

Format: **nagios\_user=<username/UID>**

Example: **nagios\_user=nagios**

This is used to set the effective user that the Nagios process should run as. After initial program startup and before starting to monitor anything, Nagios will drop its effective privileges and run as this user. You may specify either a username or a UID.

### Nagios Group

Format: **nagios\_group=<groupname/GID>**

Example: **nagios\_group=nagios**

This is used to set the effective group that the Nagios process should run as. After initial program startup and before starting to monitor anything, Nagios will drop its effective privileges and run as this group. You may specify either a groupname or a GID.

### Notifications Option

Format: **enable\_notifications=<0/1>**

Example: **enable\_notifications=1**

This option determines whether or not Nagios will send out [notifications](#) when it initially (re)starts. If this option is disabled, Nagios will not send out notifications for any host or service. Note: If you have [state retention](#) enabled, Nagios will ignore this setting when it (re)starts and use the last known setting for this option (as stored in the [state retention file](#)), *unless* you disable the [use\\_retained\\_program\\_state](#) option. If you want to change this option when state retention is active (and the [use\\_retained\\_program\\_state](#) is enabled), you'll have to use the appropriate [external command](#) or change it via the web interface. Values are as follows:

- 0 = Disable notifications
- 1 = Enable notifications (default)

### Service Check Execution Option

Format: `execute_service_checks=<0/1>`

Example: `execute_service_checks=1`

This option determines whether or not Nagios will execute service checks when it initially (re)starts. If this option is disabled, Nagios will not actively execute any service checks and will remain in a sort of "sleep" mode (it can still accept [passive checks](#) unless you've [disabled them](#)). This option is most often used when configuring backup monitoring servers, as described in the documentation on [redundancy](#), or when setting up a [distributed](#) monitoring environment. Note: If you have [state retention](#) enabled, Nagios will ignore this setting when it (re)starts and use the last known setting for this option (as stored in the [state retention file](#)), *unless* you disable the [use\\_retained\\_program\\_state](#) option. If you want to change this option when state retention is active (and the [use\\_retained\\_program\\_state](#) is enabled), you'll have to use the appropriate [external command](#) or change it via the web interface. Values are as follows:

- 0 = Don't execute service checks
- 1 = Execute service checks (default)

### Passive Service Check Acceptance Option

Format: `accept_passive_service_checks=<0/1>`

Example: `accept_passive_service_checks=1`

This option determines whether or not Nagios will accept [passive service checks](#) when it initially (re)starts. If this option is disabled, Nagios will not accept any passive service checks. Note: If you have [state retention](#) enabled, Nagios will ignore this setting when it (re)starts and use the last known setting for this option (as stored in the [state retention file](#)), *unless* you disable the [use\\_retained\\_program\\_state](#) option. If you want to change this option when state retention is active (and the [use\\_retained\\_program\\_state](#) is enabled), you'll have to use the appropriate [external command](#) or change it via the web interface. Values are as follows:

- 0 = Don't accept passive service checks
- 1 = Accept passive service checks (default)

### Host Check Execution Option

Format: `execute_host_checks=<0/1>`

Example: `execute_host_checks=1`

This option determines whether or not Nagios will execute on-demand and regularly scheduled host checks when it initially (re)starts. If this option is disabled, Nagios will not actively execute any host checks, although it can still accept [passive host checks](#) unless you've [disabled them](#)). This option is most often used when configuring backup monitoring servers, as described in the documentation on [redundancy](#), or when setting up a [distributed](#) monitoring environment. Note: If you have [state retention](#) enabled, Nagios will ignore this setting when it (re)starts and use the last known setting for this option

(as stored in the [state retention file](#)), *unless* you disable the [use\\_retained\\_program\\_state](#) option. If you want to change this option when state retention is active (and the [use\\_retained\\_program\\_state](#) is enabled), you'll have to use the appropriate [external command](#) or change it via the web interface. Values are as follows:

- 0 = Don't execute host checks
- 1 = Execute host checks (default)

### Passive Host Check Acceptance Option

Format: **accept\_passive\_host\_checks=<0/1>**

Example: **accept\_passive\_host\_checks=1**

This option determines whether or not Nagios will accept [passive host checks](#) when it initially (re)starts. If this option is disabled, Nagios will not accept any passive host checks. Note: If you have [state retention](#) enabled, Nagios will ignore this setting when it (re)starts and use the last known setting for this option (as stored in the [state retention file](#)), *unless* you disable the [use\\_retained\\_program\\_state](#) option. If you want to change this option when state retention is active (and the [use\\_retained\\_program\\_state](#) is enabled), you'll have to use the appropriate [external command](#) or change it via the web interface. Values are as follows:

- 0 = Don't accept passive host checks
- 1 = Accept passive host checks (default)

### Event Handler Option

Format: **enable\_event\_handlers=<0/1>**

Example: **enable\_event\_handlers=1**

This option determines whether or not Nagios will run [event handlers](#) when it initially (re)starts. If this option is disabled, Nagios will not run any host or service event handlers. Note: If you have [state retention](#) enabled, Nagios will ignore this setting when it (re)starts and use the last known setting for this option (as stored in the [state retention file](#)), *unless* you disable the [use\\_retained\\_program\\_state](#) option. If you want to change this option when state retention is active (and the [use\\_retained\\_program\\_state](#) is enabled), you'll have to use the appropriate [external command](#) or change it via the web interface. Values are as follows:

- 0 = Disable event handlers
- 1 = Enable event handlers (default)

### Log Rotation Method

Format: **log\_rotation\_method=<n/h/d/w/m>**

Example: **log\_rotation\_method=d**

This is the rotation method that you would like Nagios to use for your log file. Values are as follows:

- n = None (don't rotate the log - this is the default)
- h = Hourly (rotate the log at the top of each hour)
- d = Daily (rotate the log at midnight each day)
- w = Weekly (rotate the log at midnight on Saturday)
- m = Monthly (rotate the log at midnight on the last day of the month)

### Log Archive Path

Format: **log\_archive\_path=<path>**

Example: **log\_archive\_path=/usr/local/nagios/var/archives/**

This is the directory where Nagios should place log files that have been rotated. This option is ignored if you choose to not use the [log rotation](#) functionality.

### External Command Check Option

Format: **check\_external\_commands=<0/1>**

Example: **check\_external\_commands=1**

This option determines whether or not Nagios will check the [command file](#) for commands that should be executed. This option must be enabled if you plan on using the [command CGI](#) to issue commands via the web interface. More information on external commands can be found [here](#).

- 0 = Don't check external commands (default)
- 1 = Check external commands

### External Command Check Interval

Format: **command\_check\_interval=<xxx>[s]**

Example: **command\_check\_interval=1**

If you specify a number with an "s" appended to it (i.e. 30s), this is the number of *seconds* to wait between external command checks. If you leave off the "s", this is the number of "time units" to wait between external command checks. Unless you've changed the [interval\\_length](#) value (as defined below) from the default value of 60, this number will mean minutes.

Note: By setting this value to **-1**, Nagios will check for external commands as often as possible. Each time Nagios checks for external commands it will read and process all commands present in the [command file](#) before continuing on with its other duties. More information on external commands can be found [here](#).

### External Command File

Format: **command\_file=<file\_name>**

Example: **command\_file=/usr/local/nagios/var/rw/nagios.cmd**



This is the file that Nagios will check for external commands to process. The [command CGI](#) writes commands to this file. Other third party programs can write to this file if proper file permissions have been granted as outline in [here](#). The external command file is implemented as a named pipe (FIFO), which is created when Nagios starts and removed when it shuts down. If the file exists when Nagios starts, the Nagios process will terminate with an error message. More information on external commands can be found [here](#).

### Downtime File

Format: **downtime\_file=<file\_name>**

Example: **downtime\_file=/usr/local/nagios/var/downtime.dat**

This is the file that Nagios will use for storing scheduled host and service [downtime](#) information. Comments can be viewed and added for both hosts and services through the [extended information CGI](#).

### Comment File

Format: **comment\_file=<file\_name>**

Example: **comment\_file=/usr/local/nagios/var/comment.dat**

This is the file that Nagios will use for storing service and host comments. Comments can be viewed and added for both hosts and services through the [extended information CGI](#).

### Lock File

Format: **lock\_file=<file\_name>**

Example: **lock\_file=/tmp/nagios.lock**

This option specifies the location of the lock file that Nagios should create when it runs as a daemon (when started with the -d command line argument). This file contains the process id (PID) number of the running Nagios process.

### State Retention Option

Format: **retain\_state\_information=<0/1>**

Example: **retain\_state\_information=1**

This option determines whether or not Nagios will retain state information for hosts and services between program restarts. If you enable this option, you should supply a value for the [state\\_retention\\_file](#) variable. When enabled, Nagios will save all state information for hosts and service before it shuts down (or restarts) and will read in previously saved state information when it starts up again.

- 0 = Don't retain state information
- 1 = Retain state information (default)

## State Retention File

Format: **state\_retention\_file=<file\_name>**

Example: **state\_retention\_file=/usr/local/nagios/var/retention.dat**

This is the file that Nagios will use for storing service and host state information before it shuts down. When Nagios is restarted it will use the information stored in this file for setting the initial states of services and hosts before it starts monitoring anything. This file is deleted after Nagios reads in initial state information when it (re)starts. In order to make Nagios retain state information between program restarts, you must enable the [retain\\_state\\_information](#) option.

## Automatic State Retention Update Interval

Format: **retention\_update\_interval=<minutes>**

Example: **retention\_update\_interval=60**

This setting determines how often (in minutes) that Nagios will automatically save retention data during normal operation. If you set this value to 0, Nagios will not save retention data at regular intervals, but it will still save retention data before shutting down or restarting. If you have disabled state retention (with the [retain\\_state\\_information](#) option), this option has no effect.

## Use Retained Program State Option

Format: **use\_retained\_program\_state=<0/1>**

Example: **use\_retained\_program\_state=1**

This setting determines whether or not Nagios will set various program-wide state variables based on the values saved in the retention file. Some of these program-wide state variables that are normally saved across program restarts if state retention is enabled include the [enable\\_notifications](#), [enable\\_flap\\_detection](#), [enable\\_event\\_handlers](#), [execute\\_service\\_checks](#), and [accept\\_passive\\_service\\_checks](#) options. If you do not have [state retention](#) enabled, this option has no effect.

- 0 = Don't use retained program state
- 1 = Use retained program state (default)

## Use Retained Scheduling Info Option

Format: **use\_retained\_scheduling\_info=<0/1>**

Example: **use\_retained\_scheduling\_info=1**

This setting determines whether or not Nagios will retain scheduling info (next check times) for hosts and services when it restarts. If you are adding a large number (or percentage) of hosts and services, I would recommend disabling this option when you first restart Nagios, as it can adversely skew the spread of initial checks. Otherwise you will probably want to leave it enabled.

- 0 = Don't use retained scheduling info
- 1 = Use retained scheduling info (default)

### Syslog Logging Option

Format: **use\_syslog=<0/1>**

Example: **use\_syslog=1**

This variable determines whether messages are logged to the syslog facility on your local host. Values are as follows:

- 0 = Don't use syslog facility
- 1 = Use syslog facility

### Notification Logging Option

Format: **log\_notifications=<0/1>**

Example: **log\_notifications=1**

This variable determines whether or not notification messages are logged. If you have a lot of contacts or regular service failures your log file will grow relatively quickly. Use this option to keep contact notifications from being logged.

- 0 = Don't log notifications
- 1 = Log notifications

### Service Check Retry Logging Option

Format: **log\_service\_retries=<0/1>**

Example: **log\_service\_retries=1**

This variable determines whether or not service check retries are logged. Service check retries occur when a service check results in a non-OK state, but you have configured Nagios to retry the service more than once before responding to the error. Services in this situation are considered to be in "soft" states. Logging service check retries is mostly useful when attempting to debug Nagios or test out service [event handlers](#).

- 0 = Don't log service check retries
- 1 = Log service check retries

### Host Check Retry Logging Option

Format: **log\_host\_retries=<0/1>**

Example: **log\_host\_retries=1**

This variable determines whether or not host check retries are logged. Logging host check retries is mostly useful when attempting to debug Nagios or test out host [event handlers](#).

- 0 = Don't log host check retries
- 1 = Log host check retries

### Event Handler Logging Option

Format: `log_event_handlers=<0/1>`

Example: `log_event_handlers=1`

This variable determines whether or not service and host [event handlers](#) are logged. Event handlers are optional commands that can be run whenever a service or hosts changes state. Logging event handlers is most useful when debugging Nagios or first trying out your event handler scripts.

- 0 = Don't log event handlers
- 1 = Log event handlers

### Initial States Logging Option

Format: `log_initial_states=<0/1>`

Example: `log_initial_states=1`

This variable determines whether or not Nagios will force all initial host and service states to be logged, even if they result in an OK state. Initial service and host states are normally only logged when there is a problem on the first check. Enabling this option is useful if you are using an application that scans the log file to determine long-term state statistics for services and hosts.

- 0 = Don't log initial states (default)
- 1 = Log initial states

### External Command Logging Option

Format: `log_external_commands=<0/1>`

Example: `log_external_commands=1`

This variable determines whether or not Nagios will log [external commands](#) that it receives from the [external command file](#). Note: This option does not control whether or not [passive service checks](#) (which are a type of external command) get logged. To enable or disable logging of passive checks, use the [log\\_passive\\_checks](#) option.

- 0 = Don't log external commands
- 1 = Log external commands (default)

### Passive Check Logging Option

Format: **log\_passive\_checks=<0/1>**

Example: **log\_passive\_checks=1**

This variable determines whether or not Nagios will log [passive host and service checks](#) that it receives from the [external command file](#). If you are setting up a [distributed monitoring environment](#) or plan on handling a large number of passive checks on a regular basis, you may wish to disable this option so your log file doesn't get too large.

- 0 = Don't log passive checks
- 1 = Log passive checks (default)

### Global Host Event Handler Option

Format: **global\_host\_event\_handler=<command>**

Example: **global\_host\_event\_handler=log-host-event-to-db**

This option allows you to specify a host event handler command that is to be run for every host state change. The global event handler is executed immediately prior to the event handler that you have optionally specified in each host definition. The *command* argument is the short name of a command that you define in your [object configuration file](#). The maximum amount of time that this command can run is controlled by the [event\\_handler\\_timeout](#) option. More information on event handlers can be found [here](#).

### Global Service Event Handler Option

Format: **global\_service\_event\_handler=<command>**

Example: **global\_service\_event\_handler=log-service-event-to-db**

This option allows you to specify a service event handler command that is to be run for every service state change. The global event handler is executed immediately prior to the event handler that you have optionally specified in each service definition. The *command* argument is the short name of a command that you define in your [object configuration file](#). The maximum amount of time that this command can run is controlled by the [event\\_handler\\_timeout](#) option. More information on event handlers can be found [here](#).

### Inter-Check Sleep Time

Format: **sleep\_time=<seconds>**

Example: **sleep\_time=1**

This is the number of seconds that Nagios will sleep before checking to see if the next service or host check in the scheduling queue should be executed. Note that Nagios will only sleep after it "catches up" with queued service checks that have fallen behind.

### Service Inter-Check Delay Method

Format: **service\_inter\_check\_delay\_method=<n/d/s/x.xx>**

Example: **service\_inter\_check\_delay\_method=s**

This option allows you to control how service checks are initially "spread out" in the event queue. Using a "smart" delay calculation (the default) will cause Nagios to calculate an average check interval and spread initial checks of all services out over that interval, thereby helping to eliminate CPU load spikes. Using no delay is generally *not* recommended unless you are testing the [service check parallelization](#) functionality. Using no delay will cause all service checks to be scheduled for execution at the same time. This means that you will generally have large CPU spikes when the services are all executed in parallel. More information on how to estimate how the inter-check delay affects service check scheduling can be found [here](#). Values are as follows:

- n = Don't use any delay - schedule all service checks to run immediately (i.e. at the same time!)
- d = Use a "dumb" delay of 1 second between service checks
- s = Use a "smart" delay calculation to spread service checks out evenly (default)
- x.xx = Use a user-supplied inter-check delay of x.xx seconds

### Maximum Service Check Spread

Format: **max\_service\_check\_spread=<minutes>**

Example: **max\_service\_check\_spread=30**

This option determines the maximum number of minutes from when Nagios starts that all services (that are scheduled to be regularly checked) are checked. This option will automatically adjust the [service inter-check delay](#) (if necessary) to ensure that the initial checks of all services occur within the timeframe you specify. In general, this option will not have an affect on service check scheduling if scheduling information is being retained using the [use\\_retained\\_scheduling\\_info](#) option. Default value is 30 (minutes).

### Service Interleave Factor

Format: **service\_interleave\_factor=<s|x>**

Example: **service\_interleave\_factor=s**

This variable determines how service checks are interleaved. Interleaving allows for a more even distribution of service checks, reduced load on *remote* hosts, and faster overall detection of host problems. With the introduction of service check [parallelization](#), remote hosts could get bombarded with checks if interleaving was not implemented. This could cause the service checks to fail or return incorrect results if the remote host was overloaded with processing other service check requests. Setting this value to 1 is equivalent to not interleaving the service checks (this is how versions of Nagios previous to 0.0.5 worked). Set this value to s (smart) for automatic calculation of the interleave factor unless you have a specific reason to change it. The best way to understand how interleaving works is to watch the [status CGI](#) (detailed view) when Nagios is just starting. You should see that the service check results are spread out as they begin to appear. More information on how interleaving works can be found [here](#).

- $x$  = A number greater than or equal to 1 that specifies the interleave factor to use. An interleave factor of 1 is equivalent to not interleaving the service checks.
- $s$  = Use a "smart" interleave factor calculation (default)

### Maximum Concurrent Service Checks

Format: `max_concurrent_checks=<max_checks>`

Example: `max_concurrent_checks=20`

This option allows you to specify the maximum number of service checks that can be run in [parallel](#) at any given time. Specifying a value of 1 for this variable essentially prevents any service checks from being parallelized. Specifying a value of 0 (the default) does not place any restrictions on the number of concurrent checks. You'll have to modify this value based on the system resources you have available on the machine that runs Nagios, as it directly affects the maximum load that will be imposed on the system (processor utilization, memory, etc.). More information on how to estimate how many concurrent checks you should allow can be found [here](#).

### Service Reaper Frequency

Format: `service_reaper_frequency=<frequency_in_seconds>`

Example: `service_reaper_frequency=10`

This option allows you to control the frequency *in seconds* of service "reaper" events. "Reaper" events process the results from [parallelized service checks](#) that have finished executing. These events constitute the core of the monitoring logic in Nagios.

### Host Inter-Check Delay Method

Format: `host_inter_check_delay_method=<n/d/s/x.xx>`

Example: `host_inter_check_delay_method=s`

This option allows you to control how host checks *that are scheduled to be checked on a regular basis* are initially "spread out" in the event queue. Using a "smart" delay calculation (the default) will cause Nagios to calculate an average check interval and spread initial checks of all hosts out over that interval, thereby helping to eliminate CPU load spikes. Using no delay is generally *not* recommended. Using no delay will cause all host checks to be scheduled for execution at the same time. More information on how to estimate how the inter-check delay affects host check scheduling can be found [here](#). Values are as follows:

- $n$  = Don't use any delay - schedule all host checks to run immediately (i.e. at the same time!)
- $d$  = Use a "dumb" delay of 1 second between host checks
- $s$  = Use a "smart" delay calculation to spread host checks out evenly (default)
- $x.xx$  = Use a user-supplied inter-check delay of  $x.xx$  seconds

### Maximum Host Check Spread

Format: `max_host_check_spread=<minutes>`

Example: `max_host_check_spread=30`

This option determines the maximum number of minutes from when Nagios starts that all hosts (that are scheduled to be regularly checked) are checked. This option will automatically adjust the [host inter-check delay](#) (if necessary) to ensure that the initial checks of all hosts occur within the timeframe you specify. In general, this option will not have an affect on host check scheduling if scheduling information is being retained using the [use\\_retained\\_scheduling\\_info](#) option. Default value is 30 (minutes).

### Timing Interval Length

Format: `interval_length=<seconds>`

Example: `interval_length=60`

This is the number of seconds per "unit interval" used for timing in the scheduling queue, re-notifications, etc. "Units intervals" are used in the object configuration file to determine how often to run a service check, how often of re-notify a contact, etc.

**Important:** The default value for this is set to 60, which means that a "unit value" of 1 in the object configuration file will mean 60 seconds (1 minute). I have not really tested other values for this variable, so proceed at your own risk if you decide to do so!

### Auto-Rescheduling Option

Format: `auto_reschedule_checks=<0/1>`

Example: `auto_reschedule_checks=1`

This option determines whether or not Nagios will attempt to automatically reschedule active host and service checks to "smooth" them out over time. This can help to balance the load on the monitoring server, as it will attempt to keep the time between consecutive checks consistent, at the expense of executing checks on a more rigid schedule.

**WARNING:** THIS IS AN EXPERIMENTAL FEATURE AND MAY BE REMOVED IN FUTURE VERSIONS. ENABLING THIS OPTION CAN DEGRADE PERFORMANCE - RATHER THAN INCREASE IT - IF USED IMPROPERLY!

### Auto-Rescheduling Interval

Format: `auto_rescheduling_interval=<seconds>`

Example: `auto_rescheduling_interval=30`

This option determines how often (in seconds) Nagios will attempt to automatically reschedule checks. This option only has an effect if the [auto\\_reschedule\\_checks](#) option is enabled. Default is 30 seconds.



**WARNING:** THIS IS AN EXPERIMENTAL FEATURE AND MAY BE REMOVED IN FUTURE VERSIONS. ENABLING THE AUTO-RESCHEDULING OPTION CAN DEGRADE PERFORMANCE - RATHER THAN INCREASE IT - IF USED IMPROPERLY!

### Auto-Rescheduling Window

Format: `auto_rescheduling_window=<seconds>`

Example: `auto_rescheduling_window=180`

This option determines the "window" of time (in seconds) that Nagios will look at when automatically rescheduling checks. Only host and service checks that occur in the next X seconds (determined by this variable) will be rescheduled. This option only has an effect if the [auto\\_reschedule\\_checks](#) option is enabled. Default is 180 seconds (3 minutes).

**WARNING:** THIS IS AN EXPERIMENTAL FEATURE AND MAY BE REMOVED IN FUTURE VERSIONS. ENABLING THE AUTO-RESCHEDULING OPTION CAN DEGRADE PERFORMANCE - RATHER THAN INCREASE IT - IF USED IMPROPERLY!

### Aggressive Host Checking Option

Format: `use_aggressive_host_checking=<0/1>`

Example: `use_aggressive_host_checking=0`

Nagios tries to be smart about how and when it checks the status of hosts. In general, disabling this option will allow Nagios to make some smarter decisions and check hosts a bit faster. Enabling this option will increase the amount of time required to check hosts, but may improve reliability a bit. Unless you have problems with Nagios not recognizing that a host recovered, I would suggest **not** enabling this option.

- 0 = Don't use aggressive host checking (default)
- 1 = Use aggressive host checking

### Flap Detection Option

Format: `enable_flap_detection=<0/1>`

Example: `enable_flap_detection=0`

This option determines whether or not Nagios will try and detect hosts and services that are "flapping". Flapping occurs when a host or service changes between states too frequently, resulting in a barrage of notifications being sent out. When Nagios detects that a host or service is flapping, it will temporarily suppress notifications for that host/service until it stops flapping. Flap detection is very experimental at this point, so use this feature with caution! More information on how flap detection and handling works can be found [here](#). Note: If you have [state retention](#) enabled, Nagios will ignore this setting when it (re)starts and use the last known setting for this option (as stored in the [state retention file](#)), *unless* you disable the [use\\_retained\\_program\\_state](#) option. If you want to change this option when state retention is active (and the [use\\_retained\\_program\\_state](#) is enabled), you'll have to use the appropriate [external command](#) or change it via the web interface.

- 0 = Don't enable flap detection (default)
- 1 = Enable flap detection

### Low Service Flap Threshold

Format: **low\_service\_flap\_threshold=<percent>**

Example: **low\_service\_flap\_threshold=25.0**

This option is used to set the low threshold for detection of service flapping. For more information on how flap detection and handling works (and how this option affects things) read [this](#).

### High Service Flap Threshold

Format: **high\_service\_flap\_threshold=<percent>**

Example: **high\_service\_flap\_threshold=50.0**

This option is used to set the low threshold for detection of service flapping. For more information on how flap detection and handling works (and how this option affects things) read [this](#).

### Low Host Flap Threshold

Format: **low\_host\_flap\_threshold=<percent>**

Example: **low\_host\_flap\_threshold=25.0**

This option is used to set the low threshold for detection of host flapping. For more information on how flap detection and handling works (and how this option affects things) read [this](#).

### High Host Flap Threshold

Format: **high\_host\_flap\_threshold=<percent>**

Example: **high\_host\_flap\_threshold=50.0**

This option is used to set the low threshold for detection of host flapping. For more information on how flap detection and handling works (and how this option affects things) read [this](#).

### Soft Service Dependencies Option

Format: **soft\_state\_dependencies=<0/1>**

Example: **soft\_state\_dependencies=0**

This option determines whether or not Nagios will use soft service state information when checking [service dependencies](#). Normally Nagios will only use the latest hard service state when checking dependencies. If you want it to use the latest state (regardless of whether its a soft or hard [state type](#)), enable this option.

- 0 = Don't use soft service state dependencies (default)
- 1 = Use soft service state dependencies

## Service Check Timeout

Format: **service\_check\_timeout=<seconds>**

Example: **service\_check\_timeout=60**

This is the maximum number of seconds that Nagios will allow service checks to run. If checks exceed this limit, they are killed and a CRITICAL state is returned. A timeout error will also be logged.

There is often widespread confusion as to what this option really does. It is meant to be used as a last ditch mechanism to kill off plugins which are misbehaving and not exiting in a timely manner. It should be set to something high (like 60 seconds or more), so that each service check normally finishes executing within this time limit. If a service check runs longer than this limit, Nagios will kill it off thinking it is a runaway processes.

## Host Check Timeout

Format: **host\_check\_timeout=<seconds>**

Example: **host\_check\_timeout=60**

This is the maximum number of seconds that Nagios will allow host checks to run. If checks exceed this limit, they are killed and a CRITICAL state is returned and the host will be assumed to be DOWN. A timeout error will also be logged.

There is often widespread confusion as to what this option really does. It is meant to be used as a last ditch mechanism to kill off plugins which are misbehaving and not exiting in a timely manner. It should be set to something high (like 60 seconds or more), so that each host check normally finishes executing within this time limit. If a host check runs longer than this limit, Nagios will kill it off thinking it is a runaway processes.

## Event Handler Timeout

Format: **event\_handler\_timeout=<seconds>**

Example: **event\_handler\_timeout=60**

This is the maximum number of seconds that Nagios will allow [event handlers](#) to be run. If an event handler exceeds this time limit it will be killed and a warning will be logged.

There is often widespread confusion as to what this option really does. It is meant to be used as a last ditch mechanism to kill off commands which are misbehaving and not exiting in a timely manner. It should be set to something high (like 60 seconds or more), so that each event handler command normally finishes executing within this time limit. If an event handler runs longer than this limit, Nagios will kill it off thinking it is a runaway processes.

## Notification Timeout

Format: **notification\_timeout=<seconds>**

Example: **notification\_timeout=60**

This is the maximum number of seconds that Nagios will allow notification commands to be run. If a notification command exceeds this time limit it will be killed and a warning will be logged.

There is often widespread confusion as to what this option really does. It is meant to be used as a last ditch mechanism to kill off commands which are misbehaving and not exiting in a timely manner. It should be set to something high (like 60 seconds or more), so that each notification command finishes executing within this time limit. If a notification command runs longer than this limit, Nagios will kill it off thinking it is a runaway processes.

## Obsessive Compulsive Service Processor Timeout

Format: **ocsp\_timeout=<seconds>**

Example: **ocsp\_timeout=5**

This is the maximum number of seconds that Nagios will allow an [obsessive compulsive service processor command](#) to be run. If a command exceeds this time limit it will be killed and a warning will be logged.

## Obsessive Compulsive Host Processor Timeout

Format: **ochp\_timeout=<seconds>**

Example: **ochp\_timeout=5**

This is the maximum number of seconds that Nagios will allow an [obsessive compulsive host processor command](#) to be run. If a command exceeds this time limit it will be killed and a warning will be logged.

## Performance Data Processor Command Timeout

Format: **perfdata\_timeout=<seconds>**

Example: **perfdata\_timeout=5**

This is the maximum number of seconds that Nagios will allow a [host performance data processor command](#) or [service performance data processor command](#) to be run. If a command exceeds this time limit it will be killed and a warning will be logged.

## Obsess Over Services Option

Format: **obsess\_over\_services=<0/1>**

Example: **obsess\_over\_services=1**

This value determines whether or not Nagios will "obsess" over service checks results and run the [obsessive compulsive service processor command](#) you define. I know - funny name, but it was all I could think of. This option is useful for performing [distributed monitoring](#). If you're not doing distributed monitoring, don't enable this option.

- 0 = Don't obsess over services (default)
- 1 = Obsess over services

### Obsessive Compulsive Service Processor Command

Format: **ocsp\_command=<command>**

Example: **ocsp\_command=obsessive\_service\_handler**

This option allows you to specify a command to be run after *every* service check, which can be useful in [distributed monitoring](#). This command is executed after any [event handler](#) or [notification](#) commands. The *command* argument is the short name of a [command definition](#) that you define in your object configuration file. The maximum amount of time that this command can run is controlled by the [ocsp\\_timeout](#) option. More information on distributed monitoring can be found [here](#). This command is only executed if the [obsess\\_over\\_services](#) option is enabled globally and if the *obsess\_over\_service* directive in the [service definition](#) is enabled.

### Obsess Over Hosts Option

Format: **obsess\_over\_hosts=<0/1>**

Example: **obsess\_over\_hosts=1**

This value determines whether or not Nagios will "obsess" over host checks results and run the [obsessive compulsive host processor command](#) you define. I know - funny name, but it was all I could think of. This option is useful for performing [distributed monitoring](#). If you're not doing distributed monitoring, don't enable this option.

- 0 = Don't obsess over hosts (default)
- 1 = Obsess over hosts

### Obsessive Compulsive Host Processor Command

Format: **ochp\_command=<command>**

Example: **ochp\_command=obsessive\_host\_handler**

This option allows you to specify a command to be run after *every* host check, which can be useful in [distributed monitoring](#). This command is executed after any [event handler](#) or [notification](#) commands. The *command* argument is the short name of a [command definition](#) that you define in your object configuration file. The maximum amount of time that this command can run is controlled by the [ochp\\_timeout](#) option. More information on distributed monitoring can be found [here](#). This command is only executed if the [obsess\\_over\\_hosts](#) option is enabled globally and if the *obsess\_over\_host* directive in the [host definition](#) is enabled.

## Performance Data Processing Option

Format: **process\_performance\_data=<0/1>**

Example: **process\_performance\_data=1**

This value determines whether or not Nagios will process host and service check [performance data](#).

- 0 = Don't process performance data (default)
- 1 = Process performance data

## Host Performance Data Processing Command

Format: **host\_perfdata\_command=<command>**

Example: **host\_perfdata\_command=process-host-perfdata**

This option allows you to specify a command to be run after *every* host check to process host [performance data](#) that may be returned from the check. The *command* argument is the short name of a [command definition](#) that you define in your object configuration file. This command is only executed if the [process\\_performance\\_data](#) option is enabled globally and if the *process\_perf\_data* directive in the [host definition](#) is enabled.

## Service Performance Data Processing Command

Format: **service\_perfdata\_command=<command>**

Example: **service\_perfdata\_command=process-service-perfdata**

This option allows you to specify a command to be run after *every* service check to process service [performance data](#) that may be returned from the check. The *command* argument is the short name of a [command definition](#) that you define in your object configuration file. This command is only executed if the [process\\_performance\\_data](#) option is enabled globally and if the *process\_perf\_data* directive in the [service definition](#) is enabled.

## Host Performance Data File

Format: **host\_perfdata\_file=<file\_name>**

Example: **host\_perfdata\_file=/usr/local/nagios/var/host-perfdata.dat**

This option allows you to specify a file to which host [performance data](#) will be written after every host check. Data will be written to the performance file as specified by the [host\\_perfdata\\_file\\_template](#) option. Performance data is only written to this file if the [process\\_performance\\_data](#) option is enabled globally and if the *process\_perf\_data* directive in the [host definition](#) is enabled.

## Service Performance Data File

Format: **service\_perfdata\_file=<file\_name>**

Example: **service\_perfdata\_file=/usr/local/nagios/var/service-perfdata.dat**

This option allows you to specify a file to which service [performance data](#) will be written after every service check. Data will be written to the performance file as specified by the [service\\_perfdata\\_file\\_template](#) option. Performance data is only written to this file if the [process\\_performance\\_data](#) option is enabled globally and if the *process\_perf\_data* directive in the [service definition](#) is enabled.

### Host Performance Data File Template

Format: **host\_perfdata\_file\_template=<template>**

Example: **host\_perfdata\_file\_template=[HOSTPERFDATA]\t\$TIMETS\t\$HOSTNAMES\t\$HOSTEXECUTIONTIMES\t\$HOSTOUTPUTS\t\$HOSTPERFDATA\$**

This option determines what (and how) data is written to the [host performance data file](#). The template may contain [macros](#), special characters (\t for tab, \r for carriage return, \n for newline) and plain text. A newline is automatically added after each write to the performance data file.

### Service Performance Data File Template

Format: **service\_perfdata\_file\_template=<template>**

Example: **service\_perfdata\_file\_template=[SERVICEPERFDATA]\t\$TIMETS\t\$HOSTNAMES\t\$SERVICEDESCS\t\$SERVICEEXECUTIONTIMES\t\$SERVICELATENCY\t\$SERVICEOUTPUTS\t\$SERVICEPERFDATA\$**

This option determines what (and how) data is written to the [service performance data file](#). The template may contain [macros](#), special characters (\t for tab, \r for carriage return, \n for newline) and plain text. A newline is automatically added after each write to the performance data file.

### Host Performance Data File Mode

Format: **host\_perfdata\_file\_mode=<mode>**

Example: **host\_perfdata\_file\_mode=a**

This option determines whether the [host performance data file](#) is opened in write or append mode. Unless the file is a named pipe, you will probably want to use the default mode of append.

- a = Open file in append mode (default)
- w = Open file in write mode

### Service Performance Data File Mode

Format: **service\_perfdata\_file\_mode=<mode>**

Example: **service\_perfdata\_file\_mode=a**

This option determines whether the [service performance data file](#) is opened in write or append mode. Unless the file is a named pipe, you will probably want to use the default mode of append.

- a = Open file in append mode (default)
- w = Open file in write mode

### Host Performance Data File Processing Interval

Format: **host\_perfdata\_file\_processing\_interval=<seconds>**

Example: **host\_perfdata\_file\_processing\_interval=0**

This option allows you to specify the interval (in seconds) at which the [host performance data file](#) is processed using the [host performance data file processing command](#). A value of 0 indicates that the performance data file should not be processed at regular intervals.

### Service Performance Data File Processing Interval

Format: **service\_perfdata\_file\_processing\_interval=<seconds>**

Example: **service\_perfdata\_file\_processing\_interval=0**

This option allows you to specify the interval (in seconds) at which the [service performance data file](#) is processed using the [service performance data file processing command](#). A value of 0 indicates that the performance data file should not be processed at regular intervals.

### Host Performance Data File Processing Command

Format: **host\_perfdata\_file\_processing\_command=<command>**

Example: **host\_perfdata\_file\_processing\_command=process-host-perfdata-file**

This option allows you to specify the command that should be executed to process the [host performance data file](#). The *command* argument is the short name of a [command definition](#) that you define in your object configuration file. The interval at which this command is executed is determined by the [host\\_perfdata\\_file\\_processing\\_interval](#) directive.

### Service Performance Data File Processing Command

Format: **service\_perfdata\_file\_processing\_command=<command>**

Example: **service\_perfdata\_file\_processing\_command=process-service-perfdata-file**

This option allows you to specify the command that should be executed to process the [service performance data file](#). The *command* argument is the short name of a [command definition](#) that you define in your object configuration file. The interval at which this command is executed is determined by the [service\\_perfdata\\_file\\_processing\\_interval](#) directive.

### Orphaned Service Check Option



Format: **check\_for\_orphaned\_services=<0/1>**

Example: **check\_for\_orphaned\_services=1**

This option allows you to enable or disable checks for orphaned service checks. Orphaned service checks are checks which have been executed and have been removed from the event queue, but have not had any results reported in a long time. Since no results have come back in for the service, it is not rescheduled in the event queue. This can cause service checks to stop being executed. Normally it is very rare for this to happen - it might happen if an external user or process killed off the process that was being used to execute a service check. If this option is enabled and Nagios finds that results for a particular service check have not come back, it will log an error message and reschedule the service check. If you start seeing service checks that never seem to get rescheduled, enable this option and see if you notice any log messages about orphaned services.

- 0 = Don't check for orphaned service checks
- 1 = Check for orphaned service checks (default)

### Service Freshness Checking Option

Format: **check\_service\_freshness=<0/1>**

Example: **check\_service\_freshness=0**

This option determines whether or not Nagios will periodically check the "freshness" of service checks. Enabling this option is useful for helping to ensure that [passive service checks](#) are received in a timely manner. More information on freshness checking can be found [here](#).

- 0 = Don't check service freshness
- 1 = Check service freshness (default)

### Service Freshness Check Interval

Format: **service\_freshness\_check\_interval=<seconds>**

Example: **service\_freshness\_check\_interval=60**

This setting determines how often (in seconds) Nagios will periodically check the "freshness" of service check results. If you have disabled service freshness checking (with the [check\\_service\\_freshness option](#)), [this option has no effect](#). More information on freshness checking can be found [here](#).

### Host Freshness Checking Option

Format: **check\_host\_freshness=<0/1>**

Example: **check\_host\_freshness=0**

This option determines whether or not Nagios will periodically check the "freshness" of host checks. Enabling this option is useful for helping to ensure that [passive host checks](#) are received in a timely manner. More information on freshness checking can be found [here](#).

- 0 = Don't check host freshness
- 1 = Check host freshness (default)

## Host Freshness Check Interval

Format: **host\_freshness\_check\_interval=<seconds>**

Example: **host\_freshness\_check\_interval=60**

This setting determines how often (in seconds) Nagios will periodically check the "freshness" of host check results. If you have disabled host freshness checking (with the [check\\_host\\_freshness](#) option), this option has no effect. [More information on freshness checking can be found here.](#)

## Date Format

Format: **date\_format=<option>**

Example: **date\_format=us**

This option allows you to specify what kind of date/time format Nagios should use in the web interface and date/time [macros](#). Possible options (along with example output) include:

| Option         | Output Format       | Sample Output       |
|----------------|---------------------|---------------------|
| us             | MM/DD/YYYY HH:MM:SS | 06/30/2002 03:15:00 |
| euro           | DD/MM/YYYY HH:MM:SS | 30/06/2002 03:15:00 |
| iso8601        | YYYY-MM-DD HH:MM:SS | 2002-06-30 03:15:00 |
| strict-iso8601 | YYYY-MM-DDTHH:MM:SS | 2002-06-30T03:15:00 |

## Illegal Object Name Characters

Format: **illegal\_object\_name\_chars=<chars...>**

Example: **illegal\_object\_name\_chars='~!\$%^&\*"'<>?,()=**

This option allows you to specify illegal characters that cannot be used in host names, service descriptions, or names of other object types. Nagios will allow you to use most characters in object definitions, but I recommend not using the characters shown in the example above. Doing may give you problems in the web interface, notification commands, etc.

## Illegal Macro Output Characters

Format: **illegal\_macro\_output\_chars=<chars...>**

Example: **illegal\_macro\_output\_chars='~\$^&'"<>**

This option allows you to specify illegal characters that should be stripped from [macros](#) before being used in notifications, event handlers, and other commands. This DOES NOT affect macros used in service or host check commands. You can choose to not strip out the characters shown in the example above, but I recommend you do not do this. Some of these characters are interpreted by the shell (i.e. the backtick) and can lead to security problems. The following macros are stripped of the characters you specify:

**\$HOSTOUTPUT\$, \$HOSTPERFDATA\$, \$HOSTACKAUTHOR\$, \$HOSTACKCOMMENT\$,  
\$SERVICEOUTPUT\$, \$SERVICEPERFDATA\$, \$SERVICEACKAUTHOR\$, and  
\$SERVICEACKCOMMENT\$**

### Regular Expression Matching Option

Format: **use\_regex\_matching=<0/1>**

Example: **use\_regex\_matching=0**

This option determines whether or not various directives in your [object definitions](#) will be processed as regular expressions. More information on how this works can be found [here](#).

- 0 = Don't use regular expression matching (default)
- 1 = Use regular expression matching

### True Regular Expression Matching Option

Format: **use\_true\_regex\_matching=<0/1>**

Example: **use\_true\_regex\_matching=0**

If you've enabled regular expression matching of various object directives using the [use\\_regex\\_matching](#) option, this option will determine when object directives are treated as regular expressions. If this option is disabled (the default), directives will only be treated as regular expressions if they contain a \* or ? wildcard character. If this option is enabled, all appropriate directives will be treated as regular expression - be careful when enabling this! More information on how this works can be found [here](#).

- 0 = Don't use true regular expression matching (default)
- 1 = Use true regular expression matching

### Administrator Email Address

Format: **admin\_email=<email\_address>**

Example: **admin\_email=root@localhost.localdomain**

This is the email address for the administrator of the local machine (i.e. the one that Nagios is running on). This value can be used in notification commands by using the **\$ADMINEMAIL\$** [macro](#).

### Administrator Pager

Format: **admin\_pager=<pager\_number\_or\_pager\_email\_gateway>**

Example: **admin\_pager=pageroot@localhost.localdomain**

This is the pager number (or pager email gateway) for the administrator of the local machine (i.e. the one that Nagios is running on). The pager number/address can be used in notification commands by using the **\$ADMINPAGER\$** [macro](#).

---

# Object Definitions

---

## What is Object Data?

Object data is simply a generic term I use to describe various data definitions you need in order to monitor anything. Types of object definitions include:

- Services
- Service Groups
- Hosts
- Host Groups
- Contacts
- Contact Groups
- Commands
- Time Periods
- Service Escalations
- Service Dependencies
- Host Escalations
- Host Dependencies
- Extended Host Information
- Extended Service Information

## Where Is Object Data Defined?

Object data is defined in one or more configuration files that you specify using the [cfg\\_file](#) and/or [cfg\\_dir](#) directives in the [main configuration file](#). You can include multiple object configuration files and/or directories by using multiple *cfg\_file* and/or *cfg\_dir* directives.

## How Is Object Data Defined?

Object definitions are defined in a template format. [Click here](#) for more information on defining object data using this method.

---

# CGI Configuration File Options

---

## Notes

When creating and/or editing configuration files, keep the following in mind:

1. Lines that start with a '#' character are taken to be comments and are not processed
2. Variables names must begin at the start of the line - no white space is allowed before the name
3. Variable names are case-sensitive

## Sample Configuration

A sample CGI configuration file is created when you run the configure script - you can find the sample config file in the *sample-config/* subdirectory of the Nagios distribution.

## Config File Location

By default, Nagios expects the CGI configuration file to be named **cgi.cfg** and located in the config file directory along with the [main config file](#). If you need to change the name of the file or its location, you can configure Apache to pass an environment variable named NAGIOS\_CGI\_CONFIG (which points to the correct location) to the CGIs. See the Apache documentation for information on how to do this.

## Index

[Main configuration file location](#)  
[Physical HTML path](#)  
[URL HTML path](#)  
[Authentication usage](#)  
[Default user name](#)  
[System/process information access](#)  
[System/process command access](#)  
[Configuration information access](#)  
[Global host information access](#)  
[Global host command access](#)  
[Global service information access](#)  
[Global service command access](#)  
[Statusmap CGI background image](#)  
[Default statusmap layout method](#)  
[Statuswrl CGI include world](#)  
[Default statuswrl layout method](#)  
[CGI refresh rate](#)  
[Audio alerts](#)  
[Ping syntax](#)

### Main Configuration File Location

Format: **main\_config\_file=<file\_name>**

Example: **main\_config\_file=/usr/local/nagios/etc/nagios.cfg**

This specifies the location of your [main configuration file](#). The CGIs need to know where to find this file in order to get information about configuration information, current host and service status, etc.

### Physical HTML Path

Format: **physical\_html\_path=<path>**

Example: **physical\_html\_path=/usr/local/nagios/share**

This is the *physical* path where the HTML files for Nagios are kept on your workstation or server. Nagios assumes that the documentation and images files (used by the CGIs) are stored in subdirectories called *docs/* and *images/*, respectively.

### URL HTML Path

Format: **url\_html\_path=<path>**

Example: **url\_html\_path=/nagios**

If, when accessing Nagios via a web browser, you point to an URL like **http://www.myhost.com/nagios**, this value should be */nagios*. Basically, its the path portion of the URL that is used to access the Nagios HTML pages.

### Authentication Usage

Format: **use\_authentication=<0/1>**

Example: **use\_authentication=1**

This option controls whether or not the CGIs will use the authentication and authorization functionality when determining what information and commands users have access to. I would strongly suggest that you use the authentication functionality for the CGIs. If you decide not to use authentication, make sure to remove the [command CGI](#) to prevent unauthorized users from issuing commands to Nagios. The CGI will not issue commands to Nagios if authentication is disabled, but I would suggest removing it altogether just to be on the safe side. More information on how to setup authentication and configure authorization for the CGIs can be found [here](#).

- 0 = Don't use authentication functionality
- 1 = Use authentication and authorization functionality (default)

### Default User Name

Format: **default\_user\_name=<username>**

Example: **default\_user\_name=guest**

Setting this variable will define a default username that can access the CGIs. This allows people within a secure domain (i.e., behind a firewall) to access the CGIs without necessarily having to authenticate to the web server. You may want to use this to avoid having to use basic authentication if you are not using a secure server, as basic authentication transmits passwords in clear text over the Internet.

**Important:** Do *not* define a default username unless you are running a secure web server and are sure that everyone who has access to the CGIs has been authenticated in some manner! If you define this variable, anyone who has not authenticated to the web server will inherit all rights you assign to this user!

### System/Process Information Access

Format: `authorized_for_system_information=<user1>,<user2>,<user3>,...<usern>`

Example: `authorized_for_system_information=nagiosadmin,theboss`

This is a comma-delimited list of names of *authenticated users* who can view system/process information in the [extended information CGI](#). Users in this list are *not* automatically authorized to issue system/process commands. If you want users to be able to issue system/process commands as well, you must add them to the [authorized\\_for\\_system\\_commands](#) variable. More information on how to setup authentication and configure authorization for the CGIs can be found [here](#).

### System/Process Command Access

Format: `authorized_for_system_commands=<user1>,<user2>,<user3>,...<usern>`

Example: `authorized_for_system_commands=nagiosadmin`

This is a comma-delimited list of names of *authenticated users* who can issue system/process commands via the [command CGI](#). Users in this list are *not* automatically authorized to view system/process information. If you want users to be able to view system/process information as well, you must add them to the [authorized\\_for\\_system\\_information](#) variable. More information on how to setup authentication and configure authorization for the CGIs can be found [here](#).

### Configuration Information Access

Format: `authorized_for_configuration_information=<user1>,<user2>,<user3>,...<usern>`

Example: `authorized_for_configuration_information=nagiosadmin`

This is a comma-delimited list of names of *authenticated users* who can view configuration information in the [configuration CGI](#). Users in this list can view information on all configured hosts, host groups, services, contacts, contact groups, time periods, and commands. More information on how to setup authentication and configure authorization for the CGIs can be found [here](#).

### Global Host Information Access

Format: `authorized_for_all_hosts=<user1>,<user2>,<user3>,...<usern>`

Example: `authorized_for_all_hosts=nagiosadmin,theboss`

This is a comma-delimited list of names of *authenticated users* who can view status and configuration information for all hosts. Users in this list are also automatically authorized to view information for all services. Users in this list are *not* automatically authorized to issue commands for all hosts or services. If you want users able to issue commands for all hosts and services as well, you must add them to the



`authorized_for_all_host_commands` variable. More information on how to setup authentication and configure authorization for the CGIs can be found [here](#).

### Global Host Command Access

Format: `authorized_for_all_host_commands=<user1>,<user2>,<user3>,...<usern>`

Example: `authorized_for_all_host_commands=nagiosadmin`

This is a comma-delimited list of names of *authenticated users* who can issue commands for all hosts via the [command CGI](#). Users in this list are also automatically authorized to issue commands for all services. Users in this list are *not* automatically authorized to view status or configuration information for all hosts or services. If you want users able to view status and configuration information for all hosts and services as well, you must add them to the [authorized\\_for\\_all\\_hosts](#) variable. More information on how to setup authentication and configure authorization for the CGIs can be found [here](#).

### Global Service Information Access

Format: `authorized_for_all_services=<user1>,<user2>,<user3>,...<usern>`

Example: `authorized_for_all_services=nagiosadmin,theboss`

This is a comma-delimited list of names of *authenticated users* who can view status and configuration information for all services. Users in this list are *not* automatically authorized to view information for all hosts. Users in this list are *not* automatically authorized to issue commands for all services. If you want users able to issue commands for all services as well, you must add them to the [authorized\\_for\\_all\\_service\\_commands](#) variable. More information on how to setup authentication and configure authorization for the CGIs can be found [here](#).

### Global Service Command Access

Format: `authorized_for_all_service_commands=<user1>,<user2>,<user3>,...<usern>`

Example: `authorized_for_all_service_commands=nagiosadmin`

This is a comma-delimited list of names of *authenticated users* who can issue commands for all services via the [command CGI](#). Users in this list are *not* automatically authorized to issue commands for all hosts. Users in this list are *not* automatically authorized to view status or configuration information for all hosts. If you want users able to view status and configuration information for all services as well, you must add them to the [authorized\\_for\\_all\\_services](#) variable. More information on how to setup authentication and configure authorization for the CGIs can be found [here](#).

### Statusmap CGI Background Image

Format: `statusmap_background_image=<image_file>`

Example: `statusmap_background_image=smbbackground.gd2`

This option allows you to specify an image to be used as a background in the [statusmap CGI](#) if you use the user-supplied coordinates layout method. The background image is not available in any other layout methods. It is assumed that the image resides in the HTML images path (i.e. /usr/local/nagios/share/images). This path is automatically determined by appending "/images" to the path specified by the [physical\\_html\\_path](#) directive. Note: The image file can be in GIF, JPEG, PNG, or GD2 format. However, GD2 format (preferably in uncompressed format) is recommended, as it will reduce the CPU load when the CGI generates the map image.

### Default Statusmap Layout Method

Format: **default\_statusmap\_layout=<layout\_number>**

Example: **default\_statusmap\_layout=4**

This option allows you to specify the default layout method used by the [statusmap CGI](#). Valid options are:

| <layout_number> Value | Layout Method            |
|-----------------------|--------------------------|
| 0                     | User-defined coordinates |
| 1                     | Depth layers             |
| 2                     | Collapsed tree           |
| 3                     | Balanced tree            |
| 4                     | Circular                 |
| 5                     | Circular (Marked Up)     |
| 6                     | Circular (Balloon)       |

### Statuswrl CGI Include World

Format: **statuswrl\_include=<vrml\_file>**

Example: **statuswrl\_include=myworld.wrl**

This option allows you to include your own objects in the generated VRML world. It is assumed that the file resides in the path specified by the [physical\\_html\\_path](#) directive. Note: This file must be a fully qualified VRML world (i.e. you can view it by itself in a VRML browser).

### Default Statuswrl Layout Method

Format: **default\_statuswrl\_layout=<layout\_number>**

Example: **default\_statuswrl\_layout=4**

This option allows you to specify the default layout method used by the [statuswrl CGI](#). Valid options are:

| <layout_number> Value | Layout Method            |
|-----------------------|--------------------------|
| 0                     | User-defined coordinates |
| 2                     | Collapsed tree           |
| 3                     | Balanced tree            |
| 4                     | Circular                 |

### CGI Refresh Rate

Format: **refresh\_rate=<rate\_in\_seconds>**

Example: **refresh\_rate=90**

This option allows you to specify the number of seconds between page refreshes for the [status](#), [statusmap](#), and [extinfo](#) CGIs.

### Audio Alerts

Formats: **host\_unreachable\_sound=<sound\_file>**  
**host\_down\_sound=<sound\_file>**  
**service\_critical\_sound=<sound\_file>**  
**service\_warning\_sound=<sound\_file>**  
**service\_unknown\_sound=<sound\_file>**

Examples: **host\_unreachable\_sound=hostu.wav**  
**host\_down\_sound=hostd.wav**  
**service\_critical\_sound=critical.wav**  
**service\_warning\_sound=warning.wav**  
**service\_unknown\_sound=unknown.wav**

These options allow you to specify an audio file that should be played in your browser if there are problems when you are viewing the [status CGI](#). If there are problems, the audio file for the most critical type of problem will be played. The most critical type of problem is on or more unreachable hosts, while the least critical is one or more services in an unknown state (see the order in the example above). Audio files are assumed to be in the **media/** subdirectory in your HTML directory (i.e. */usr/local/nagios/share/media*).

### Ping Syntax

Format: **ping\_syntax=<command>**

Example: **ping\_syntax=/bin/ping -n -U -c 5 \$HOSTADDRESS\$**

This option determines what syntax should be used when attempting to ping a host from the WAP interface (using the [statuswml CGI](#)). You must include the full path to the ping binary, along with all required options. The \$HOSTADDRESS\$ macro is substituted with the address of the host before the command is executed.

---

# Authentication And Authorization In The CGIs

---

## Notes

Throughout these instructions I will be assuming that you are running the [Apache](#) web server on your machine. If you are running some other web server, you will have to make some adjustments.

## Definitions

Throughout these instructions I will be using the following terms, so you should understand what they mean...

- An *authenticated user* is an someone who has authenticated to the web server with a username and password and has been granted access to the Nagios web interface.
- An *authenticated contact* is an authenticated user whose username matches the short name of a contact definition in your [object configuration file\(s\)](#).

## Index

[Setting up authenticated users](#)

[Enabling authentication/authorization functionality in the CGIs](#)

[Default permissions to CGI information](#)

[Granting additional permissions to CGI information](#)

[Authentication on secure web servers](#)

## Setting Up Authenticated Users

If you haven't done so already, you'll need to add the appropriate entries to your web server config file to enable basic authentication for the CGI and HTML portions of the Nagios web interface. Instructions for doing so can be found [here](#).

Now that you've configured your web server to require authentication for the Nagios web interface, you'll need to specify who has access. This is done by using the **htpasswd** command supplied with Apache.

Running the following command will create a new file called *htpasswd.users* in the */usr/local/nagios/etc* directory. It will also create an username/password entry for *nagiosadmin*. You will be asked to provide a password that will be used when *nagiosadmin* authenticates to the web server.

**`htpasswd -c /usr/local/nagios/etc/htpasswd.users nagiosadmin`**

Continue adding more users until you've created an account for everyone you want to access the CGIs. Use the following command to add additional users, replacing `<username>` with the actual username you want to add. Note that the **-c** option is not used, since you already created the initial file.

**`htpasswd /usr/local/nagios/etc/htpasswd.users <username>`**

Okay, so you're done with the first part of what needs to be done. At this point you should be prompted for a username and password if you point your web browser to the Nagios web interface. If you have problems getting user authentication to work at this point, read your webserver documentation for more info.

## Enabling Authentication/Authorization Functionality In The CGIs

The next thing you need to do is make sure that the CGIs are configured to use the authentication and authorization functionality in determining what information and/or commands users have access to. This is done by setting the `use_authentication` variable in the [CGI configuration file](#) to a non-zero value. Example:

**`use_authentication=1`**

Okay, you're now done with setting up basic authentication/authorization functionality in the CGIs.

## Default Permissions To CGI Information

So what default permissions do users have in the CGIs by default when the authentication/authorization functionality is enabled?

| CGI Data                          | Authenticated Contacts <sup>*</sup> | Other Authenticated Users <sup>*</sup> |
|-----------------------------------|-------------------------------------|----------------------------------------|
| Host Status Information           | Yes                                 | No                                     |
| Host Configuration Information    | Yes                                 | No                                     |
| Host History                      | Yes                                 | No                                     |
| Host Notifications                | Yes                                 | No                                     |
| Host Commands                     | Yes                                 | No                                     |
| Service Status Information        | Yes                                 | No                                     |
| Service Configuration Information | Yes                                 | No                                     |
| Service History                   | Yes                                 | No                                     |
| Service Notifications             | Yes                                 | No                                     |
| Service Commands                  | Yes                                 | No                                     |
| All Configuration Information     | No                                  | No                                     |
| System/Process Information        | No                                  | No                                     |
| System/Process Commands           | No                                  | No                                     |

*Authenticated contacts<sup>\*</sup>* are granted the following permissions for each **service** for which they are contacts (but not for services for which they are not contacts)...

- Authorization to view service status information
- Authorization to view service configuration information
- Authorization to view history and notifications for the service
- Authorization to issue service commands

*Authenticated contacts<sup>\*</sup>* are granted the following permissions for each **host** for which they are contacts (but not for hosts for which they are not contacts)...

- Authorization to view host status information
- Authorization to view host configuration information
- Authorization to view history and notifications for the host

- Authorization to issue host commands
- Authorization to view status information for all services on the host
- Authorization to view configuration information for all services on the host
- Authorization to view history and notification information for all services on the host
- Authorization to issue commands for all services on the host

It is important to note that by default **no one** is authorized for the following...

- Viewing the raw log file via the [showlog CGI](#)
- Viewing Nagios process information via the [extended information CGI](#)
- Issuing Nagios process commands via the [command CGI](#)
- Viewing host group, contact, contact group, time period, and command definitions via the [configuration CGI](#)

You will undoubtedly want to access this information, so you'll have to assign additional rights for yourself (and possibly other users) as described below...

### Granting Additional Permissions To CGI Information

You can grant *authenticated contacts* or other *authenticated users* permission to additional information in the CGIs by adding them to various authorization variables in the [CGI configuration file](#). I realize that the available options don't allow for getting really specific about particular permissions, but its better than nothing..

Additional authorization can be given to users by adding them to the following variables in the CGI configuration file...

- [authorized\\_for\\_system\\_information](#)
- [authorized\\_for\\_system\\_commands](#)
- [authorized\\_for\\_configuration\\_information](#)
- [authorized\\_for\\_all\\_hosts](#)
- [authorized\\_for\\_all\\_host\\_commands](#)
- [authorized\\_for\\_all\\_services](#)
- [authorized\\_for\\_all\\_service\\_commands](#)

### CGI Authorization Requirements

If you are confused about the authorization needed to access various information in the CGIs, read the *Authorization Requirements* section for each CGI as described [here](#).

### Authentication On Secured Web Servers

If your web server is located in a secure domain (i.e., behind a firewall) or if you are using SSL, you can define a default username that can be used to access the CGIs. This is done by defining the [default\\_user\\_name](#) option in the [CGI configuration file](#). By defining a default username that can access the CGIs, you can allow users to access the CGIs without necessarily having to authenticate to the web server.. You may want to use this to avoid having to use basic web authentication, as basic authentication transmits passwords in clear text over the Internet.

**Important:** Do *not* define a default username unless you are running a secure web server and are sure that everyone who has access to the CGIs has been authenticated in some manner! If you define this variable, anyone who has not authenticated to the web server will inherit all rights you assign to this user!





# Verifying Your Nagios Configuration

---

## Verifying The Configuration From The Command Line

Once you've entered all the necessary data into the [configuration files](#), its time to do a sanity check. Everyone make mistakes from time to time, so its best to verify what you've entered. Nagios automatically runs a "pre-flight check" before before it starts monitoring, but you also have the option of running this check manually before attempting to start Nagios. In order to do this, you must start Nagios with the **-v** command line argument as follows...

**`/usr/local/nagios/bin/nagios -v <main_config_file>`**

Note that you should be entering the path/filename of your *main* configuration file (i.e. `/usr/local/nagios/etc/nagios.cfg`) as the second argument. Nagios will read your [main configuration file](#) and all [object configuration files](#) and verify that they contain valid data.

## Relationships Verified During The Pre-Flight Check

During the "pre-flight check", Nagios verifies that you have defined the data relationships necessary for monitoring. Objects are all related and need to be setup properly in order for things to run. This is a list of the basic things that Nagios attempts to check before it will start monitoring...

1. Verify that all contacts are a member of at least one contact group.
2. Verify that all contacts specified in each contact group are valid.
3. Verify that all hosts are a member of at least one host group.
4. Verify that all hosts specified in each host group are valid.
5. Verify that all hosts have at least one service associated with them.
6. Verify that all commands used in service and host checks are valid.
7. Verify that all commands used in service and host event handlers are valid.
8. Verify that all commands used in contact service and host notifications are valid.
9. Verify that all notification time periods specified for services, hosts, and contact are valid.
10. Verify that all service check time periods specified for services are valid.

## Fixing Configuration Errors

If you've forgotten to enter some critical data or just plain screwed things up, Nagios will spit out a warning or error message that should point you to the location of the problem. Error messages generally print out the line in the configuration file that seems to be the source of the problem. On errors, Nagios will often exit the pre-flight check and return to the command prompt after printing only the first error that it has encountered. This is done so that one error does not cascade into multiple errors as the remainder of the configuration data is verified. If you get any error messages you'll need to go and edit your configuration files to remedy the problem. Warning messages can *generally* be safely ignored, since they are only recommendations and not requirements.

## Where To Go From Here

Once you've verified your configuration files and fixed any errors, you can be reasonably sure that Nagios will start monitoring the services you've specified. On to [starting Nagios](#)!

---

# Starting Nagios

---

**IMPORTANT:** Before you actually start Nagios, you'll have to make sure that you have [configured](#) it properly and [verified the config data](#)!

## Methods For Starting Nagios

There are basically four different ways you can start Nagios:

1. Manually, as a foreground process (useful for initial testing and debugging)
2. Manually, as a background process
3. Manually, as a daemon
4. Automatically at system boot

Let's examine each method briefly...

### Running Nagios Manually as a Foreground Process

If you enabled the debugging options when running the configure script (and recompiled Nagios), this would be your first choice for testing and debugging. Running Nagios as a foreground process at a shell prompt will allow you to more easily view what's going on in the monitoring and notification processes. To run Nagios as a foreground process for testing, invoke Nagios like this...

**`/usr/local/nagios/bin/nagios <main_config_file>`**

Note that you must specify the path/filename of the [main configuration file](#) (i.e. `/usr/local/nagios/etc/nagios.cfg`) on the command line.

To stop Nagios at any time, just press CTRL-C. If you've enabled the debugging options you'll probably want to redirect the output to a file for easier review later.

### Running Nagios Manually as a Background Process

To run Nagios as a background process, invoke it with an ampersand as follows...

**`/usr/local/nagios/bin/nagios <main_config_file> &`**

Note that you must specify the path/filename of the [main configuration file](#) (i.e. `/usr/local/nagios/etc/nagios.cfg`) on the command line.

### Running Nagios Manually as a Daemon

In order to run Nagios in daemon mode you must supply the **-d** switch on the command line as follows...

**`/usr/local/nagios/bin/nagios -d <main_config_file>`**

Note that you must specify the path/filename of the [main configuration file](#) (i.e. `/usr/local/nagios/etc/nagios.cfg`) on the command line.

### Running Nagios Automatically at System Boot

When you have tested Nagios and are reasonably sure that it is not going to crash, you will probably want to have it start automatically at boot time. To do this (in Linux) you will have to create a startup script in your `/etc/rc.d/init.d/` directory. You will also have to create a link to the script in the runlevel(s)

that you wish to have Nagios to start in. I'll assume that you know what I'm talking about and are able to do this.

A sample init script (named **daemon-init**) is created in the base directory of the Nagios distribution when you run the configure script. You can install the sample script to your /etc/rc.d/init.d directory using the '**make install-init**' command, as outlined in the [installation](#) instructions.

The sample init scripts are designed to work under Linux, so if you want to use them under FreeBSD, Solaris, etc. you may have to do a little hacking...

### **Stopping and Restarting Nagios**

Directions on how to stop and restart Nagios can be found [here](#).

---

# Stopping And Restarting Nagios

Once you have Nagios up and running, you may need to stop the process or reload the configuration data "on the fly". This section describes how to do just that.

**IMPORTANT:** Before you restart Nagios, make sure that you have [verified the configuration data](#) using the -v command line switch, *especially* if you have made any changes to your [config files](#). If Nagios encounters problem with one of the config files when it restarts, it will log an error and terminate.

## Stopping And Restarting With The Init Script

If you have installed the sample init script to your /etc/rc.d/init.d directory you can stop and restart Nagios easily. If you haven't, skip this section and read how to do it manually below. I'll assume that you named the init script **Nagios** in the examples below...

| Desired Action            | Command                                | Description                                                                                                                                          |
|---------------------------|----------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------|
| Stop Nagios               | <b>/etc/rc.d/init.d/nagios stop</b>    | This kills the Nagios process                                                                                                                        |
| Restart Nagios            | <b>/etc/rc.d/init.d/nagios restart</b> | This kills the current Nagios process and then starts Nagios up again                                                                                |
| Reload Configuration Data | <b>/etc/rc.d/init.d/nagios reload</b>  | Sends a SIGHUP to the Nagios process, causing it to flush its current configuration data, reread the configuration files, and start monitoring again |

Stopping, restarting, and reloading Nagios are fairly simple with an init script and I would highly recommend you use one if at all possible.

## Stopping and Restarting Nagios Manually

If you aren't using an init script to start Nagios, you'll have to do things manually. First you'll have to find the process ID that Nagios is running under and then you'll have to use the *kill* command to terminate the application or make it reload the configuration data by sending it the proper signal. Directions for doing this are outlined below...

### Finding The Nagios Process ID

First off, you will need to know the process id that Nagios is running as. To do that, just type the following command at a shell prompt:

**ps axu | grep nagios**

The output should look something like this:

```
nagios 6808 0.0 0.7 840 352 p3 S 13:44 0:00 grep nagios
nagios 11149 0.2 1.0 868 488 ? S Feb 27 6:33 /usr/local/nagios/bin/nagios nagios.cfg
```

From the program output, you will notice that Nagios was started by user **nagios** and is running as process id **11149**.

### Manually Stopping Nagios

In order to stop Nagios, use the *kill* command as follows...

**kill 11149**

You should replace **11149** with the actual process id that Nagios is running as on your machine.

### **Manually Restarting Nagios**

If you have modified the configuration data, you will want to restart Nagios and have it re-read the new configuration. If you have changed the source code and recompiled the main Nagios executable you should *not* use this method. Instead, stop Nagios by killing it (as outlined above) and restart it manually. Restarting Nagios using the method below does not actually reload Nagios - it just causes Nagios to flush its current configuration, re-read the new configuration, and start monitoring all over again. To restart Nagios, you need to send the **SIGHUP** signal to Nagios. Assuming that the process id for Nagios is **11149** (taken from the example above), use the following command:

**kill -HUP 11149**

Remember, you will need to replace **11149** with the actual process id that Nagios is running as on your machine.

---

# Nagios Plugins

---

## What Are Plugins?

Plugins are compiled executables or scripts (Perl, shell, etc.) that can be run from a command line to check the status of a host or service. Nagios uses the results from plugins to determine the current status of hosts and services on your network. No, you can't get away without using plugins - Nagios is useless without them.

## Obtaining Plugins

Plugin development for Nagios is being done at SourceForge. The Nagios plugin development project page (where the latest version of the plugins can always be found) is located at <http://sourceforge.net/projects/nagiosplug/>.

## How Do I Use Plugin X?

Documentation on how to use individual plugins is *not* supplied with the core Nagios distribution. You should refer to the latest plugin distribution for information on using plugins. Karl DeBisschop, lead plugin developer/maintainer points out the following:

All plugins that comply with minimal development guideline for this project include internal documentation. The documentation can be read executing plugin with the '-h' option ('--help' if long options are enabled). If the '-h' option does not work, that is a bug.

For example, if you want to know how the check\_http plugin works or what options it accepts, you should try executing one of the following commands:

```
./check_http --help
```

or

```
./check_http -h
```

## Command Definition Examples For Services

It is important to note that command definitions found in sample config files in the core Nagios distribution are probably *not* accurate as to command line parameters, etc when it comes to the plugins. They are simply provided as examples of how to define commands.

## Creating Custom Plugins

Creating your own plugins to perform custom host or service checks is easy. You can find information on how to write plugins at <http://sourceforge.net/projects/nagiosplug/>. The developer guidelines can be found at <http://nagiosplug.sourceforge.net/developer-guidelines.html>.

---

# Nagios Addons

---

Several "addons" are available for Nagios on the Nagios downloads page - <http://www.nagios.org/download/>.

Addons are available for:

- Managing the config files through a web interface
  - Monitoring remote hosts (\*NIX, Windows, etc.)
  - Submitting passive checks from remote hosts
  - Simplifying/extending the notification logic
  - ...and much more
-

# Determining Status and Reachability of Network Hosts

---

## Monitoring Services on Down or Unreachable Hosts

The main purpose of Nagios is to monitor services that run on or are provided by physical hosts or devices on your network. It should be obvious that if a host or device on your network goes down, all services that it offers will also go down with it. Similarly, if a host becomes unreachable, Nagios will not be able to monitor the services associated with that host.

Nagios recognizes this fact and attempts to check for such a scenario when there are problems with a service. Whenever a service check results in a non-OK status level, Nagios will attempt to check and see if the host that the service is running on is "alive". Typically this is done by pinging the host and seeing if any response is received. If the host check command returns a non-OK state, Nagios assumes that there is a problem with the host. In this situation Nagios will "silence" all potential alerts for services running on the host and just notify the appropriate contacts that the host is down or unreachable. If the host check command returns an OK state, Nagios will recognize that the host is alive and will send out an alert for the service that is misbehaving.

## Local Hosts

"Local" hosts are hosts that reside on the same network segment as the host running Nagios - no routers or firewalls lay between them. [Figure 1](#) shows an example network layout. Host A is running Nagios and monitoring all other hosts and routers depicted in the diagram. Hosts B, C, D, E and F are all considered to be "local" hosts in relation to host A.

The `<parents>` option in the host definition for a "local" host should be left blank, as local hosts have no dependencies or "parents" - that's why they're local.

## Monitoring Local Hosts

Checking hosts that are on your local network is fairly simple. Short of someone accidentally (or intentionally) unplugging the network cable from one of your hosts, there isn't too much that can go wrong as far as checking network connectivity is concerned. There are no routers or external networks between the host doing the monitoring and the other hosts on the local network.

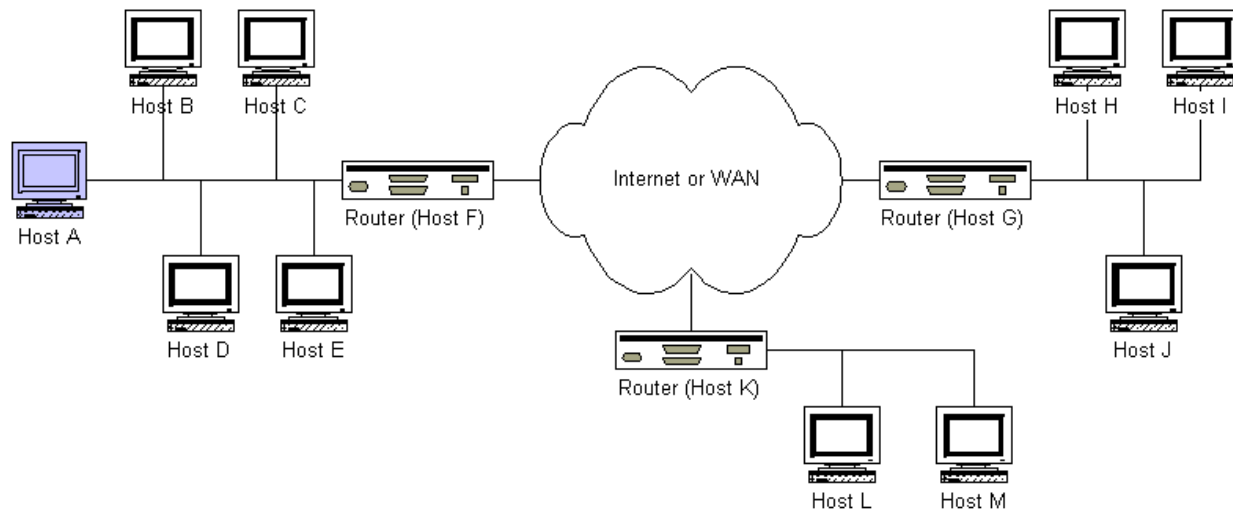
If Nagios needs to check to see if a local host is "alive" it will simply run the host check command for that host. If the command returns an OK state, Nagios assumes the host is up. If the command returns any other status level, Nagios will assume the host is down.

**Figure 1.**



## Example Network Layout

Last Modified 5/31/1999



### Remote Hosts

"Remote" hosts are hosts that reside on a different network segment than the host running Nagios. In the figure above, hosts G, H, I, J, K, L and M are all considered to be "remote" hosts in relation to host A.

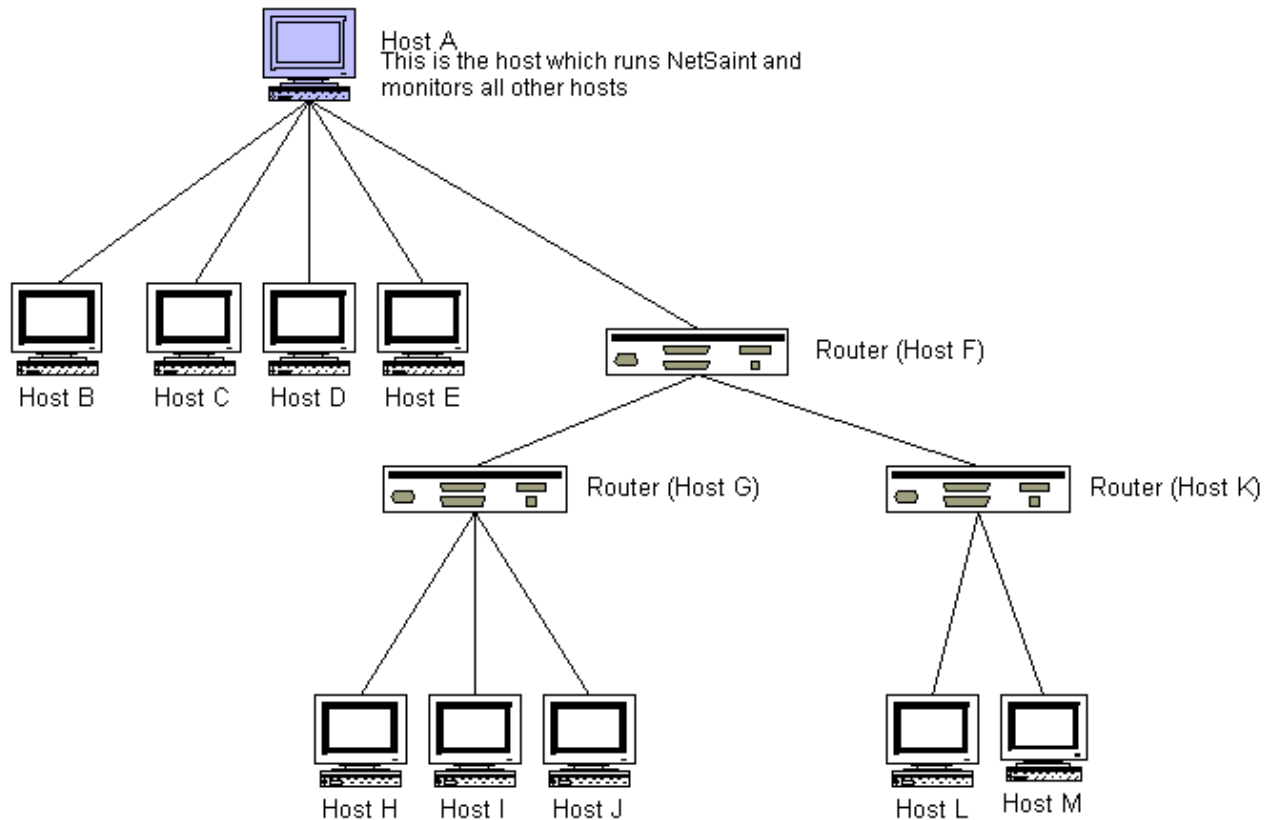
Notice that some hosts are "farther away" than others. Hosts H, I and J are one hop further away from host A than host G (the router) is. From this observation we can construct a host dependency tree as show below in [Figure 2](#). This tree diagram will help us in deciding how to configure each host in Nagios.

The **<parents>** option in the host definition for a "remote" host should be the short name(s) of the host(s) directly above it in the tree diagram (as show below). For example, the parent host for host H would be host G. The parent host for host G is host F. Host F has no parent host, since it is on the network segment as host A - it is a "local" host.

**Figure 2.**

## Network Link Heirarchy

Last Modified 5/31/1999



### Monitoring Remote Hosts

Checking the status of remote hosts is a bit more complicated than for local hosts. If Nagios cannot monitor services on a remote host, it needs to determine whether the remote host is down or whether it is unreachable. Luckily, the **<parents>** option allows Nagios to do this.

If a host check command for a remote host returns a non-OK state, Nagios will "walk" the dependency tree (as shown in the figure above) until it reaches the top (or until a parent host check results in an OK state). By doing this, Nagios is able to determine if a service problem is the result of a down host, a down network link, or just a plain old service failure.

### DOWN vs. UNREACHABLE Notification Types

I get lots of email from people asking why Nagios is sending notifications out about hosts that are unreachable. The answer is because you configured it to do that. If you want to disable UNREACHABLE notifications for hosts, modify the *notification\_options* argument of your host definitions to not include the *u* (unreachable) option.

---

# Network Outages

---

## Introduction

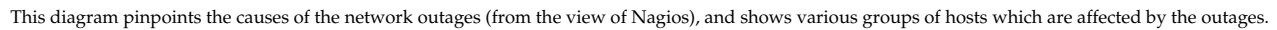
The [outages CGI](#) is designed to help pinpoint the cause of network outages. For small networks this CGI may not be particularly useful, but for larger ones it will be. Pinpointing the cause of outages will help admins to more quickly find and resolve problems which are causing the biggest impact on the network.

It should be noted that the outages CGI will not attempt to find the *exact* cause of the problem, but will rather locate the hosts on your network which seem to be causing the most problems. Delving into the problem at a deeper level is left to the user, as there are any number of things which might actually be the cause of the problem.

## Diagrams

The diagrams below help to show how the outages CGI goes about determining the cause of network outages. You can click on either image for a larger version...

This diagram will serve as the basis for our example. All hosts shows in red are either down or unreachable (from the view of Nagios). All other hosts are up.



## **Determining The Cause Of Network Outages**

So how does the outages CGI determine which hosts are the source of problems? *"Problem" hosts must be either in a DOWN or UNREACHABLE state and at least one of their immediate parent hosts must be UP.* Hosts which fit this criteria are flagged as being potential problem hosts.

In order to determine whether these flagged hosts are causing network outages, we must perform some other tests...

If *all* of the immediate child hosts of one of these flagged hosts is DOWN or UNREACHABLE *and* has no immediate parent host that is up, the flagged host is the cause of a network outage. If even one of the immediate children of a flagged host does *not* pass this test, then the flagged host is *not* the cause of a network outage.

## **Determining The Effects Of Network Outages**

Along with telling you what hosts are causing problem on your network, the outages CGI will also tell you how many hosts and services are affected by a particular problem host. How is this determined? Take a look at diagram 2 above...

From the diagram it is clear that host 1 is blocking two child hosts (in domain A). Host 2 is solely responsible for blocking only itself (domain B) and host 3 is solely responsible for blocking 7 hosts (domain C). The outage effects of the two hosts in domain D are "shared" between hosts 2 and 3, since it is unclear as to which host is actually the cause of the outage. If either host 2 or 3 was UP, then these hosts might not be blocked.

The numbers of affected hosts for each problem host are as follows (the problem host is also included in these figures):

- Host 1: 3 affected hosts
- Host 2: 3 affected hosts
- Host 3: 10 affected hosts

## **Ranking Problems Based On Severity Level**

The outages CGI will display all problem hosts, whether they are causing network outages or not. However, the CGI will tell you how many of the problem hosts (if any) are causing network outages.

In order to display the problem hosts in a somewhat useful manner, they are sorted by the severity of the effect they are having on the network. The severity level is determined by two things: The number of hosts which are affected by problem host and the number of services which are affected. Hosts hold a higher weight than services when it comes to calculating severity. The current code sets this weight ratio at 4:1 (i.e. hosts are 4 times more important than individual services).

Assuming that all hosts in diagram 2 have an equal number of services associated with them, host 3 would be ranked as the most severe problem, while hosts 1 and 2 would have the same severity level.

---

# Notifications

---

## Introduction

I've had a lot of questions as to exactly how notifications work. This will attempt to explain exactly when and how host and service notifications are sent out, as well as who receives them.

## When Do Notifications Occur?

The decision to send out notifications is made in the service check and host check logic. Host and service notifications occur in the following instances...

- When a hard state change occurs. More information on state types and hard state changes can be found [here](#).
- When a host or service remains in a hard non-OK state and the time specified by the `<notification_interval>` option in the host or service definition has passed since the last notification was sent out (for that specified host or service). If you don't like the idea of recurring notifications, set the `<notification_interval>` value to 0 - this prevents notifications from getting sent out more than once for any given problem.

## Who Gets Notified?

Each service definition has a `<contact_groups>` option that specifies what contact groups receive notifications for that particular service. Each contact group can contain one or more individual contacts. When Nagios sends out a service notification, it will notify each contact that is a member of any contact groups specified in the `<contactgroups>` option of the service definition. Nagios realizes that any given contact may be a member of more than one contact group, so it removes duplicate contact notifications before it does anything.

Each host definition has a `<contact_groups>` option that specifies what contact groups receive notifications for that particular host. When Nagios sends out a host notification, it will notify contacts that are members of all the contact groups that that should be notified for that host. Nagios removes any duplicate contacts from the notification list before it does anything.

## What Filters Must Be Passed In Order For Notifications To Be Sent?

Just because there is a need to send out a host or service notification doesn't mean that any contacts are going to get notified. There are several filters that potential notifications must pass before they are deemed worthy enough to be sent out. Even then, specific contacts may not be notified if their notification filters do not allow for the notification to be sent to them. Let's go into the filters that have to be passed in more detail...

### Program-Wide Filter:

The first filter that notifications must pass is a test of whether or not notifications are enabled on a program-wide basis. This is initially determined by the `enable_notifications` directive in the main config file, but may be changed during runtime from the web interface. If notifications are disabled on a program-wide basis, no host or service notifications can be sent out - period. If they are enabled on a program-wide basis, there are still other tests that must be passed...

### Service and Host Filters:

The first filter for host or service notifications is a check to see if the host or service is in a period of [scheduled downtime](#). If it is in a scheduled downtime, **no one gets notified**. If it isn't in a period of downtime, it gets passed on to the next filter. As a side note, notifications for services are suppressed if the host they're associated with is in a period of scheduled downtime.

The second filter for host or service notification is a check to see if the host or service is [flapping](#) (if you enabled flap detection). If the service or host is currently flapping, **no one gets notified**. Otherwise it gets passed to the next filter.

The third host or service filter that must be passed is the host- or service-specific notification options. Each service definition contains options that determine whether or not notifications can be sent out for warning states, critical states, and recoveries. Similarly, each host definition contains options that determine whether or not notifications can be sent out when the host goes down, becomes unreachable, or recovers. If the host or service notification does not pass these options, **no one gets notified**. If it does pass these options, the notification gets passed to the next filter... Note: Notifications about host or service recoveries are only sent out if a notification was sent out for the original problem. It doesn't make sense to get a recovery notification for something you never knew was a problem.

The fourth host or service filter that must be passed is the time period test. Each host and service definition has a `<notification_period>` option that specifies which time period contains valid notification times for the host or service. If the time that the notification is being made does not fall within a valid time range in the specified time period, **no one gets contacted**. If it falls within a valid time range, the notification gets passed to the next filter... Note: If the time period filter is not passed, Nagios will reschedule the next notification for the host or service (if its in a non-OK state) for the next valid time present in the time period. This helps ensure that contacts are notified of problems as soon as possible when the next valid time in time period arrives.

The last set of host or service filters is conditional upon two things: (1) a notification was already sent out about a problem with the host or service at some point in the past and (2) the host or service has remained in the same non-OK state that it was when the last notification went out. If these two criteria are met, then Nagios will check and make sure the time that has passed since the last notification went out either meets or exceeds the value specified by the `<notification_interval>` option in the host or service definition. If not enough time has passed since the last notification, **no one gets contacted**. If either enough time has passed since the last notification or the two criteria for this filter were not met, the notification will be sent out! Whether or not it actually is sent to individual contacts is up to another set of filters...

### **Contact Filters:**

At this point the notification has passed the program mode filter and all host or service filters and Nagios starts to notify [all the people it should](#). Does this mean that each contact is going to receive the notification? No! Each contact has their own set of filters that the notification must pass before they receive it. Note: Contact filters are specific to each contact and do not affect whether or not other contacts receive notifications.

The first filter that must be passed for each contact are the notification options. Each contact definition contains options that determine whether or not service notifications can be sent out for warning states, critical states, and recoveries. Each contact definition also contains options that determine whether or not host notifications can be sent out when the host goes down, becomes unreachable, or recovers. If the host or service notification does not pass these options, **the contact will not be notified**. If it does pass these options, the notification gets passed to the next filter... Note: Notifications about host or service recoveries are only sent out if a notification was sent out for the original problem. It doesn't make sense to get a recovery notification for something you never knew was a problem...

The last filter that must be passed for each contact is the time period test. Each contact definition has a `<notification_period>` option that specifies which time period contains valid notification times for the contact. If the time that the notification is being made does not fall within a valid time range in the specified time period, **the contact will not be notified**. If it falls within a valid time range, the contact gets notified!

### What Aren't Any Notification Methods Incorporated Directly Into Nagios?

I've gotten several questions about why notification methods (paging, etc.) are not directly incorporated into the Nagios code. The answer is simple - it just doesn't make much sense. The "core" of Nagios is not designed to be an all-in-one application. If service checks were embedded in Nagios' core it would be very difficult for users to add new check methods, modify existing checks, etc. Notifications work in a similar manner. There are a thousand different ways to do notifications and there are already a lot of packages out there that handle the dirty work, so why re-invent the wheel and limit yourself to a bike tire? It's much easier to let an external entity (i.e. a simple script or a full-blown messaging system) do the messy stuff. Some messaging packages that can handle notifications for pagers and cellphones are listed below in the resource section.

### Notification Type Macro

When crafting your notification commands, you need to take into account what type of notification is occurring. The `$NOTIFICATIONTYPE$` macro contains a string that identifies exactly that. The table below lists the possible values for the macro and their respective descriptions:

| Value           | Description                                                                                                                                                                                                                                                          |
|-----------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| PROBLEM         | A service or host has just entered (or is still in) a problem state. If this is a service notification, it means the service is either in a WARNING, UNKNOWN or CRITICAL state. If this is a host notification, it means the host is in a DOWN or UNREACHABLE state. |
| RECOVERY        | A service or host recovery has occurred. If this is a service notification, it means the service has just returned to an OK state. If it is a host notification, it means the host has just returned to an UP state.                                                 |
| ACKNOWLEDGEMENT | This notification is an acknowledgement notification for a host or service problem. Acknowledgement notifications are initiated via the web interface by contacts for the particular host or service.                                                                |
| FLAPPINGSTART   | The host or service has just started <a href="#">flapping</a> .                                                                                                                                                                                                      |
| FLAPPINGSTOP    | The host or service has just stopped <a href="#">flapping</a> .                                                                                                                                                                                                      |

### Helpful Resources

There are many ways you could configure Nagios to send notifications out. It's up to you to decide which method(s) you want to use. Once you do that you'll have to install any necessary software and configure notification commands in your config files before you can use them. Here are just a few possible notification methods:

- Email
- Pager
- Phone (SMS)
- WinPopup message
- Yahoo, ICQ, or MSN instant message



- Audio alerts
- etc...

Basically anything you can do from a command line can be tailored for use as a notification command.

If you're interested in sending an alphanumeric notification to your pager or cellphone via email, you may be find the following information useful. Here are a few links to various messaging service providers' websites that contain information on how to send alphanumeric messages to pagers and phones...

- [Cingular](#)
- [PageNet](#)
- [SprintPCS](#) (SMS phones)

If you're looking for an alternative to using email for sending messages to your pager or cellphone, check out these packages. They could be used in conjunction with Nagios to send out a notification via a modem when a problem arises. That way you don't have to rely on email to send notifications out (remember, email may \*not\* work if there are network problems). I haven't actually tried these packages myself, but others have reported success using them...

- [Gnokii](#) (SMS software for contacting Nokia phones via GSM network)
- [QuickPage](#) (alphanumeric pager software)
- [Sendpage](#) (paging software)
- [SMS Client](#) (command line utility for sending messages to pagers and mobile phones)

If you want to try out a non-traditional method of notification, you might want to mess around with audio alerts. If you want to have audio alerts played on the monitoring server (with synthesized speech), check out [Festival](#). If you'd rather leave the monitoring box alone and have audio alerts played on another box, check out the [Network Audio System \(NAS\)](#) and [rplay](#) projects.

Lastly, there in an area in the contrib downloads section on the [Nagios homepage](#) for notification scripts that have been contributed by users. You might find these scripts useful, as they take care of a lot of the dirty work needed to send out alphanumeric notifications...

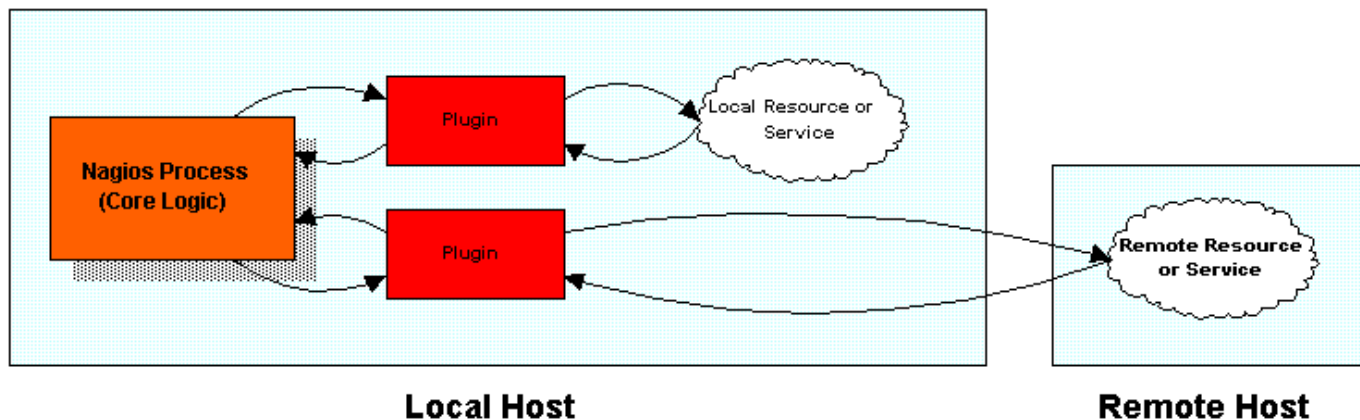
---

# Plugin Theory

## Introduction

Unlike many other monitoring tools, Nagios does not include any internal mechanisms for checking the status of services, hosts, etc. Instead, Nagios relies on external programs (called plugins) to do the all the dirty work. Nagios will execute a plugin whenever there is a need to check a service or host that is being monitored. The plugin does *something* (notice the very general term) to perform the check and then simply returns the results to Nagios. Nagios will process the results that it receives from the plugin and take any necessary actions (running [event handlers](#), sending out [notifications](#), etc).

The image below show how plugins are separated fromt the core program logic in Nagios. Nagios executes the plugins which then check local or remote resources or services of some type. When the plugins have finished checking the resource or service, they simply pass the results of the check back to Nagios for processing. A more complex diagram on how plugins work can be found in the documentation on [passive service checks](#).



## The Upside

The good thing about the plugin architecture is that you can monitor just about anything you can think of. If you can automate the process of checking something, you can monitor it with Nagios. There are already a lot of plugins that have been created in order to monitor basic resources such as processor load, disk usage, ping rates, etc. If you want to monitor something else, take a look at the documentation on [writing plugins](#) and roll your own. Its simple!

## The Downside

The only real downside to the plugin architecture is the fact that Nagios has absolutely no idea what it is that you're monitoring. You could be monitoring network traffic statistics, data error rates, room temperate, CPU voltage, fan speed, processor load, disk space, or the ability of your super-fantastic toaster to properly brown your bread in the morning... As such, Nagios cannot produce graphs of changes to the exact values of resources you're monitoring over time. It can only track changes in the *state* of those resources. Only the plugins themselves know exactly what they're monitoring and how to perform checks. However, plugins can return optional [performance data](#) along with status information. This performance data can then be passed on to external applications which could produce graphs of service-specific information (i.e. disk space usage, processor load, etc.). More information on performance data can be found [here](#).

## **Using Plugins For Service Checks**

The correlation between plugins and service checks should be fairly obvious. When Nagios needs to check the status of a particular service that you have defined, it will execute the plugin you specified in the `<check_command>` argument of the service definition. The plugin will check the status of the service or resource you specify and return the results to Nagios.

## **Using Plugins For Host Checks**

Using plugins to check the status of hosts may be a bit more difficult to understand. In each host definition you use the `<host_check_command>` argument to specify a plugin that should be executed to check the status of the host. Host checks are not performed on a regular basis - they are executed only as needed, usually when there are problems with one or more services that are associated with the host.

Host checks can use the same plugins as service checks. The only real difference between the two types of checks is in the interpretation of the plugin results. If a plugin that is used for a host check results in a non-OK status, Nagios will believe that the host is down.

In most situations, you'll want to use a plugin which checks to see if the host can be pinged, as this is the most common method of telling whether or not a host is up. However, if you were monitoring some kind of super-fantastic toaster, you might want to use a plugin that would check to see if the heating elements turned on when the handle was pushed down. That would give a decent indication as to whether or not the toaster was "alive".

---

# Service Check Scheduling

---

## Index

[Introduction](#)  
[Configuration options](#)  
[Initial scheduling](#)  
[Inter-check delay](#)  
[Service interleaving](#)  
[Max concurrent service checks](#)  
[Time restraints](#)  
[Normal scheduling](#)  
[Scheduling during problems](#)  
[Host checks](#)  
[Scheduling delays](#)  
[Scheduling example](#)  
[Service definition options that affect scheduling](#)

## Introduction

I've gotten a lot of questions regarding how service checks are scheduled in certain situations, along with how the scheduling differs from when the checks are actually executed and their results are processed. I'll try to go into a little more detail on how this all works...

## Configuration Options

Before we begin, there are several configuration options that affect how service checks are scheduled, executed, and processed. For starters, each service definition contains three options that determine when and how each specific service check is scheduled and executed. Those three options include:

- *normal\_check\_interval*
- *retry\_check\_interval*
- *check\_period*

There are also four configuration options in the [main configuration file](#) that affect service checks. These include:

- *service\_inter\_check\_delay\_method*
- *service\_interleave\_factor*
- *max\_concurrent\_checks*
- *service\_reaper\_frequency*

We'll go into more detail on how all these options affect service check scheduling as we progress. First off, let's see how services are initially scheduled when Nagios first starts or restarts...

## Initial Scheduling

When Nagios (re)starts, it will attempt to schedule the initial check of all services in a manner that will minimize the load imposed on the local and remote hosts. This is done by spacing the initial service checks out, as well as interleaving them. The spacing of service checks (also known as the inter-check delay) is used to minimize/equalize the load on the local host running Nagios and the interleaving is used to minimize/equalize load imposed on remote hosts. Both the inter-check delay and interleave functions are discussed below.

Even though service checks are initially scheduled to balance the load on both the local and remote hosts, things will eventually give in to the ensuing chaos and be a bit random. Reasons for this include the fact that services are not all checked at the same interval, some services take longer to execute than others, host and/or service problems can alter the timing of one or more service checks, etc. At least we try to get things off to a good start. Hopefully the initial scheduling will keep the load on the local and remote hosts fairly balanced as time goes by...

**Note:** If you want to view the initial service check scheduling information, start Nagios using the `-s` command line option. Doing so will display basic scheduling information (inter-check delay, interleave factor, first and last service check time, etc) and will create a new status log that shows the exact time that all services are initially scheduled. Because this option will overwrite the status log, you should not use it when another copy of Nagios is running. Nagios does *not* start monitoring anything when this argument is used.

### Inter-Check Delay

As mentioned before, Nagios attempts to equalize the load placed on the machine that is running Nagios by equally spacing out initial service checks. The spacing between consecutive service checks is called the inter-check delay. By giving a value to the `service_inter_check_delay_method` variable in the main config file, you can modify how this delay is calculated. I will discuss how the "smart" calculation works, as this is the setting you will want to use for normal operation.

When using the "smart" setting of the `service_inter_check_delay_method` variable, Nagios will calculate an inter-check delay value by using the following calculation:

*inter-check delay = (average check interval for all services) / (total number of services)*

*Let's take an example. Say you have 1,000 services that each have a normal check interval of 5 minutes (obviously some services are going to be checked at different intervals, but let's look at an easy case...). The total check interval time for all services is 5,000 (1,000 \* 5). That means that the average check interval for each service is 5 minutes (5,000 / 1,000). Give that information, we realize that (on average) we need to re-check 1,000 services every 5 minutes. This means that we should use an inter-check delay of 0.005 minutes (0.3 seconds) when spacing out the initial service checks. By spacing each service check out by 0.3 seconds, we can somewhat guarantee that Nagios is scheduling and/or executing 3 new service checks every second. By spacing the checks out evenly over time like this, we can hope that the load on the local server that is running Nagios remains somewhat balanced.*

### Service Interleaving

*As discussed above, the inter-check delay helps to equalize the load that Nagios imposes on the local host. What about remote hosts? Is it necessary to equalize load on remote hosts? Why? Yes, it is important and yes, Nagios can help out with this. Equalizing load on remote hosts is especially important with the advent of [service check parallelization](#). If you monitor a large number of services on a remote host and the checks were not spread out, the remote host might think that it was the victim of a SYN attack if there were a lot of open connections on the same port. Plus, attempting to equalize the load on hosts is just a nice thing to do...*

*By giving a value to the `service_interleave_factor` variable in the main config file, you can modify how the interleave factor is calculated. I will discuss how the "smart" calculation works, as this will probably be the setting you will want to use for normal operation. You can, however, use a pre-set interleave factor instead of having Nagios calculate one for you. Also of note, if you use an interleave factor of 1, service check interleaving is basically disabled.*

*When using the "smart" setting of the `service_interleave_factor` variable, Nagios will calculate an interleave factor by using the following calculation:*

$\text{interleave factor} = \text{ceil}(\text{total number of services} / \text{total number of hosts})$

Let's take an example. Say you have a total of 1,000 services and 150 hosts that you monitor. Nagios would calculate the interleave factor to be 7. This means that when Nagios schedules initial service checks it will schedule the first one it finds, skip the next 6, schedule the next one, and so on... This process will keep repeating until all service checks have been scheduled. Since services are sorted (and thus scheduled) by the name of the host they are associated with, this will help with minimizing/equalizing the load placed upon remote hosts.

The images below depict how service checks are scheduled when they are not interleaved (service\_interleave\_factor=1) and when they are interleaved with the service\_interleave\_factor variable equal to 4.

Non-Interleaved Checks:

| Host | Service               | Status  | Last Check | Duration     | Attempt | Service Information                                 |
|------|-----------------------|---------|------------|--------------|---------|-----------------------------------------------------|
| 01   | Apache/Proxmox/Load   | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:25 2001 |
|      | Qemu.C-File/Space     | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:27 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:29 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:31 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:33 2001 |
| 02   | Qemu                  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:35 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:37 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:39 2001 |
|      | Qemu.C-File/Space     | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:41 2001 |
|      | Physical/Meminfo/Load | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:43 2001 |
| 03   | Qemu                  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:45 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:47 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:49 2001 |
|      | Qemu.C-File/Space     | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:51 2001 |
|      | Physical/Meminfo/Load | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:53 2001 |
| 04   | Qemu                  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:55 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:57 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:04:59 2001 |
|      | Qemu.C-File/Space     | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:05:01 2001 |
|      | Physical/Meminfo/Load | PENDING | NA         | 04:00:00:15m | 0/5     | Service check scheduled for Wed Aug 1 12:05:03 2001 |

Interleaved Checks:

| Host | Service               | Status  | Last Check | Duration    | Attempt | Service Information                                 |
|------|-----------------------|---------|------------|-------------|---------|-----------------------------------------------------|
| 01   | Apache/Proxmox/Load   | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:03:00 2001 |
|      | Qemu.C-File/Space     | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:04:14 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:05:28 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:06:42 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:07:56 2001 |
| 02   | Qemu                  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:03:15 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:04:29 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:05:43 2001 |
|      | Qemu.C-File/Space     | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:06:57 2001 |
|      | Physical/Meminfo/Load | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:08:11 2001 |
| 03   | Qemu                  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:03:30 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:04:44 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:05:58 2001 |
|      | Qemu.C-File/Space     | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:07:12 2001 |
|      | Physical/Meminfo/Load | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:08:26 2001 |
| 04   | Qemu                  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:03:45 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:04:59 2001 |
|      | Physical/Meminfo/Use  | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:06:13 2001 |
|      | Qemu.C-File/Space     | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:07:27 2001 |
|      | Physical/Meminfo/Load | PENDING | NA         | 04:00:00:3m | 0/5     | Service check scheduled for Wed Aug 1 12:08:41 2001 |

Notice how checks are scheduled consecutively, in the order they appear in the list. This is due to the fact that the interleave factor is equal to 1.

Notice how every 4th check in this list is consecutively scheduled. This is due to the fact that the interleave factor is equal to 4.

| Host | Service              | Status  | Last Check          | Duration     | Attempt | Service Information                                 |
|------|----------------------|---------|---------------------|--------------|---------|-----------------------------------------------------|
| 01   | Apache/Proxmox/Load  | OK      | 08-01-2001 12:14:41 | 04:00:00:25s | 1/5     | Load ok: 1.5m avg/2.3% 5.5m avg/1.1% 15.0m avg/0.5% |
|      | Qemu.C-File/Space    | OK      | 08-01-2001 12:14:46 | 04:00:00:15s | 1/5     | Disk space ok: 717 MB (24%) of 2896 MB used         |
|      | Physical/Meminfo/Use | OK      | 08-01-2001 12:14:48 | 04:00:00:22s | 1/5     | Physical memory ok: 204.4 MB (24%) of 854 MB used   |
|      | Physical/Meminfo/Use | OK      | 08-01-2001 12:14:47 | 04:00:00:20s | 1/5     | Physical memory ok: 105.9 MB (12%) of 854 MB used   |
|      | Qemu                 | OK      | 08-01-2001 12:14:46 | 04:00:00:15s | 1/5     | PMIO ok: Packed load=0.0% RTA=0.00 ms               |
| 02   | Qemu                 | OK      | 08-01-2001 12:14:51 | 04:00:00:16s | 1/5     | PMIO ok: Packed load=0.0% RTA=0.00 ms               |
|      | Physical/Meminfo/Use | OK      | 08-01-2001 12:14:53 | 04:00:00:14s | 1/5     | Physical memory ok: 74.8 MB (28%) of 265.6 MB used  |
|      | Physical/Meminfo/Use | OK      | 08-01-2001 12:14:55 | 04:00:00:12s | 1/5     | Physical memory ok: 38.9 MB (11%) of 357.7 MB used  |
|      | Qemu.C-File/Space    | OK      | 08-01-2001 12:14:50 | 04:00:00:20s | 1/5     | Disk space ok: 2095 MB (25%) of 8396 MB used        |
|      | Apache/Proxmox/Load  | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:14:58 2001 |
| 03   | Qemu                 | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:01 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:02 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:04 2001 |
|      | Qemu.C-File/Space    | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:06 2001 |
|      | Apache/Proxmox/Load  | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:08 2001 |
| 04   | Qemu                 | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:10 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:12 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:14 2001 |
|      | Qemu.C-File/Space    | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:16 2001 |
|      | Apache/Proxmox/Load  | PENDING | NA                  | 04:00:00:20s | 0/5     | Service check scheduled for Wed Aug 1 12:15:18 2001 |

Notice how services are checked consecutively, in the order they appear in the list. Compare this to the example in which services are checked with a service interleave factor greater than 1.

| Host | Service              | Status  | Last Check          | Duration     | Attempt | Service Information                                 |
|------|----------------------|---------|---------------------|--------------|---------|-----------------------------------------------------|
| 01   | Apache/Proxmox/Load  | OK      | 08-01-2001 12:03:00 | 04:00:19.72s | 1/5     | Load ok: 1.5m avg/2.3% 5.5m avg/1.1% 15.0m avg/0.5% |
|      | Qemu.C-File/Space    | OK      | 08-01-2001 12:03:00 | 04:00:19.72s | 1/5     | Disk space ok: 717 MB (24%) of 2896 MB used         |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:03:13 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:03:21 2001 |
|      | Qemu                 | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:03:29 2001 |
| 02   | Qemu                 | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:03:38 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:03:46 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:03:54 2001 |
|      | Qemu.C-File/Space    | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:04:02 2001 |
|      | Apache/Proxmox/Load  | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:04:10 2001 |
| 03   | Qemu                 | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:04:19 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:04:27 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:04:35 2001 |
|      | Qemu.C-File/Space    | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:04:43 2001 |
|      | Apache/Proxmox/Load  | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:04:51 2001 |
| 04   | Qemu                 | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:05:00 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:05:08 2001 |
|      | Physical/Meminfo/Use | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:05:16 2001 |
|      | Qemu.C-File/Space    | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:05:24 2001 |
|      | Apache/Proxmox/Load  | PENDING | NA                  | 04:00:39.56s | 0/5     | Service check scheduled for Wed Aug 1 12:05:32 2001 |

Since the service checks are interleaved, the load reported on remote servers for the checks at any given time is somewhat balanced.

## Maximum Concurrent Service Checks

In order to prevent Nagios from consuming all of your CPU resources, you can restrict the maximum number of concurrent service checks that can be running at any given time. This is controlled by using the `max_concurrent_checks` option in the main config file.

The good thing about this setting is that you can regulate Nagios' CPU usage. The down side is that service checks may fall behind if this value is set too low. When it comes time to execute a service check, Nagios will make sure that no more than  $x$  service checks are either being executed or waiting to have their results processed (where  $x$  is the number of checks you specified for the `max_concurrent_checks` option). If that limit has been reached, Nagios will postpone the execution of any pending checks until some of the previous checks have completed. So how does one determine a reasonable value for the `max_concurrent_checks` option?

First off, you need to know the following things...

- The inter-check delay that Nagios uses to initially schedule service checks (use the `-s` command line argument to check this)
- The frequency (in seconds) of service reaper events, as specified by the `service_reaper_frequency` variable in the



main config file.

- A general idea of the average time that service checks actually take to execute (most plugins timeout after 10 seconds, so the average is probably going to be lower)

Next, use the following calculation to determine a reasonable value for the maximum number of concurrent checks that are allowed...

$$\text{max. concurrent checks} = \text{ceil}(\text{max}(\text{service reaper frequency}, \text{average check execution time}) / \text{inter-check delay})$$

The calculated number should provide a reasonable starting point for the `max_concurrent_checks` variable. You may have to increase this value a bit if service checks are still falling behind schedule or decrease it if Nagios is hogging too much CPU time.

Let's say you are monitoring 875 services, each with an average check interval of 2 minutes. That means that your inter-check delay is going to be 0.137 seconds. If you set the service reaper frequency to be 10 seconds, you can calculate a rough value for the max. number of concurrent checks as follows (I'll assume that the average execution time for service checks is less than 10 seconds) ...

$$\text{max. concurrent checks} = \text{ceil}(10 / 0.137)$$

In this case, the calculated value is going to be 73. This makes sense because (on average) Nagios are going to be executing just over 7 new service checks per second and it only processes service check results every 10 seconds. That means at given time there will be a just over 70 service checks that are either being executed or waiting to have their results processed. In this case, I would probably recommend bumping the max. concurrent checks value up to 80, since there will be delays when Nagios processes service check results and does its other work. Obviously, you're going to have test and tweak things a bit to get everything running smoothly on your system, but hopefully this provided some general guidelines...

### **Time Restraints**

The `check_period` option determines the [time period](#) during which Nagios can run checks of the service. Regardless of what status a particular service is in, if the time that it is actually executed is not a valid time within the time period that has been specified, the check will not be executed. Instead, Nagios will reschedule the service check for the next valid time in the time period. If the check can be run (e.g. the time is valid within the time period), the service check is executed.

**Note:** Even though a service check may not be able to be executed at a given time, Nagios may still schedule it to be run at that time. This is most likely to happen during the initial scheduling of services, although it may happen in other instances as well. This does not mean that Nagios will execute the check! When it comes time to actually execute a service check, Nagios will verify that the check can be run at the current time. If it cannot, Nagios will not execute the service check, but will instead just reschedule it for a later time. Don't let this one throw you confuse you! The scheduling and execution of service checks are two distinctly different (although related) things.

### **Normal Scheduling**

In an ideal world you wouldn't have network problems. But if that were the case, you wouldn't need a network monitoring tool. Anyway, when things are running smoothly and a service is in an OK state, we'll call that "normal". Service checks are normally scheduled at the frequency specified by the `check_interval` option. That's it. Simple, huh?

### **Scheduling During Problems**

So what happens when there are problems with a service? Well, one of the things that happens is the service check scheduling changes. If you've configured the `max_attempts` option of the service definition to be something greater than 1, Nagios will recheck the service before deciding that a real problem exists. While the service is being rechecked (up to `max_attempts` times) it is considered to be in a "soft" state (as described [here](#)) and the service

checks are rescheduled at a frequency determined by the `retry_interval` option.

If Nagios rechecks the service `max_attempts` times and it is still in a non-OK state, Nagios will put the service into a "hard" state, send out notifications to contacts (if applicable), and start rescheduling future checks of the service at a frequency determined by the `check_interval` option.

As always, there are exceptions to the rules. When a service check results in a non-OK state, Nagios will check the host that the service is associated with to determine whether or not is up (see the note [below](#) for info on how this is done). If the host is not up (i.e. it is either down or unreachable), Nagios will immediately put the service into a hard non-OK state and it will reset the current attempt number to 1. Since the service is in a hard non-OK state, the service check will be rescheduled at the normal frequency specified by the `check_interval` option instead of the `retry_interval` option.

## Host Checks

Unlike service checks, host checks are not scheduled on a regular basis. Instead they are run on demand, as Nagios sees a need. This is a common question asked by users, so it needs to be clarified.

One instance where Nagios checks the status of a host is when a service check results in a non-OK status. Nagios checks the host to decide whether or not the host is up, down, or unreachable. If the first host check returns a non-OK state, Nagios will keep pounding out checks of the host until either (a) the maximum number of host checks (specified by the `max_attempts` option in the host definition) is reached or (b) a host check results in an OK state.

Also of note - when Nagios is check the status of a host, it holds off on doing anything else (executing new service checks, processing other service check results, etc). This can slow things down a bit and cause pending service checks to be delayed for a while, but it is necessary to determine the status of the host before Nagios can take any further action on the service(s) that are having problems.

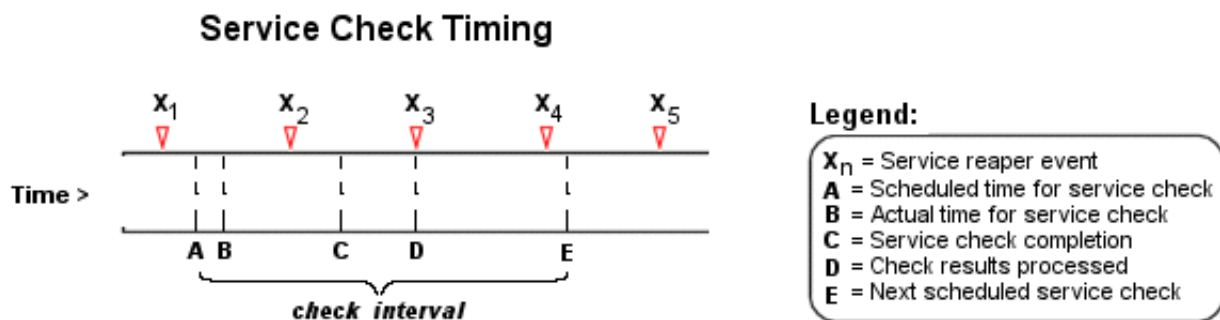
## Scheduling Delays

It should be noted that service check scheduling and execution is done on a best effort basis. Individual service checks are considered to be low priority events in Nagios, so they can get delayed if high priority events need to be executed. Examples of high priority events include log file rotations, external command checks, and service reaper events. Additionally, host checks will slow down the execution and processing of service checks.

## Scheduling Example

The scheduling of service checks, their execution, and the processing of their results can be a bit difficult to understand, so let's look at a simple example. Look at the diagram below - I'll refer to it as I explain how things are done.

Image 5.





First off, the  $X_n$  events are service reaper events that are scheduled at a frequency specified by the [service\\_reaper\\_frequency](#) option in the main config file. Service reaper events do the work of gathering and processing service check results. They serve as the core logic for Nagios, kicking off host checks, event handlers and notifications as necessary.

For the example here, a service has been scheduled to be executed at time **A**. However, Nagios got behind in its event queue, so the check was not actually executed until time **B**. The service check finished executing at time **C**, so the difference between points **C** and **B** is the actual amount of time that the check was running.

The results of the service check are not processed immediately after the check is done executing. Instead, the results are saved for later processing by a service reaper event. The next service reaper event occurs at time **D**, so that is approximately the time that the results are processed (the actual time may be later than **D** since other service check results may be processed before this one).

At the time that the service reaper event processes the service check results, it will reschedule the next service check and place it into Nagios' event queue. We'll assume that the service check resulted in an OK status, so the next check at time **E** is scheduled after the originally scheduled check time by a length of time specified by the `check_interval` option. Note that the service is not rescheduled based off the time that it was actually executed! There is one exception to this (isn't there always?) - if the time that the service check is actually executed (point **B**) occurs after the next service check time (point **E**), Nagios will compensate by adjusting the next check time. This is done to ensure that Nagios doesn't go nuts trying to keep up with service checks if it comes under heavy load. Besides, what's the point of scheduling something in the past...?

### **Service Definition Options That Affect Scheduling**

Each service definition contains a `normal_check_interval` and `retry_check_interval` option. Hopefully this will clarify what these two options do, how they relate to the `max_check_attempts` option in the service definition, and how they affect the scheduling of the service.

First off, the `normal_check_interval` option is the interval at which the service is checked under "normal" circumstances. "Normal" circumstances mean whenever the service is in an OK state or when its in a [hard](#) non-OK state.

When a service first changes from an OK state to a non-OK state, Nagios gives you the ability to temporarily slow down or speed up the interval at which subsequent checks of that service will occur. When the service first changes state, Nagios will perform up to `max_check_attempts-1` retries of the service check before it decides its a real problem. While the service is being retried, it is scheduled according to the `retry_check_interval` option, which might be faster or slower than the normal `normal_check_interval` option. While the service is being rechecked (up to `max_check_attempts-1` times), the service is in a [soft state](#). If the service is rechecked `max_check_attempts-1` times and it is still in a non-OK state, the service turns into a [hard state](#) and is subsequently rescheduled at the normal rate specified by the `check_interval` option.

On a side note, if you specify a value of 1 for the `max_check_attempts` option, the service will not ever be checked at the interval specified by the `retry_check_interval` option. Instead, it immediately turns into a [hard state](#) and is subsequently rescheduled at the rate specified by the `normal_check_interval` option.

---

# State Types

---

## Introduction

The current state of services and hosts is determined by two components: the status of the service or host (i.e. OK, WARNING, UP, DOWN, etc.) and the *type* of state it is in. There are two state types in Nagios - "soft" states and "hard" states. State types are a crucial part of Nagios' monitoring logic. They are used to determine when [event handlers](#) are executed and when notifications are sent out.

## Service and Host Check Retries

In order to prevent false alarms, Nagios allows you to define how many times a service or host check will be retried before the service or host is considered to have a real problem. The maximum number of retries before a service or host check is considered to have a real problem is controlled by the `<max_check_attempts>` option in the service and host definitions, respectively. Depending on what attempt a service or host check is currently on determines what type of state it is in. There are a few exceptions to this in the service monitoring logic, but we'll ignore those for now. Let's take a look at the different service state types...

## Soft States

Soft states occur for services and hosts in the following situations...

- When a service or host check results in a non-OK state and it has not yet been (re)checked the number of times specified by the `<max_check_attempts>` option in the service or host definition. Let's call this a soft error state...
- When a service or host recovers from a soft error state. This is considered to be a soft recovery.

## Soft State Events

What happens when a service or host is in a soft error state or experiences a soft recovery?

- The soft error or recovery is logged if you enabled the [log\\_service\\_retries](#) or [log\\_host\\_retries](#) options in the main configuration file.
- [Event handlers](#) are executed (if you defined any) to handle the soft error or recovery for the service or host. (Before any event handler is executed, the `$HOSTSTATETYPE$` or `$SERVICESTATETYPE$` [macro](#) is set to "SOFT").
- Nagios does *not* send out notifications to any contacts because there is (or was) no "real" problem with the service or host.

As can be seen, the only important thing that really happens during a soft state is the execution of event handlers. Using event handlers can be particularly useful if you want to try and proactively fix a problem before it turns into a hard state. More information on event handlers can be found [here](#).

## Hard States

Hard states occur for *services* in the following situations (hard host states are discussed later)...

- When a service check results in a non-OK state and it has been (re)checked the number of times specified by the `<max_check_attempts>` option in the service definition. This is a hard error state.
- When a service recovers from a hard error state. This is considered to be a hard recovery.
- When a service check results in a non-OK state and its corresponding host is either DOWN or UNREACHABLE. This is an exception to the general monitoring logic, but makes perfect sense. If the host isn't up why should we try and recheck the service?

Hard states occur for *hosts* in the following situations...

- When a host check results in a non-OK state and it has been (re)checked the number of times specified by the `<max_check_attempts>` option in the host definition. This is a hard error state.
- When a host recovers from a hard error state. This is considered to be a hard recovery.

### **Hard State Changes**

Before I discuss what happens when a host or service is in a hard state, you need to know about hard state changes. Hard state changes occur when a service or host...

- changes from a hard OK state to a hard non-OK state
- changes from a hard non-OK state to a hard OK-state
- changes from a hard non-OK state of some kind to a hard non-OK state of another kind (i.e. from a hard WARNING state to a hard UNKNOWN state)

### **Hard State Events**

What happens when a service or host is in a hard error state or experiences a hard recovery? Well, that depends on whether or not a hard state change (as described above) has occurred.

If a hard state change has occurred *and* the service or host is in a non-OK state the following things will occur..

- The hard service or host problem is logged.
- [Event handlers](#) are executed (if you defined any) to handle the hard problem for the service or host. (Before any event handler is executed, the `$HOSTSTATETYPE$` or `$SERVICESTATETYPE$` macro is set to "**HARD**").
- Contacts will be notified of the service or host problem (if the [notification logic](#) allows it).

If a hard state change has occurred *and* the service or host is in an OK state the following things will occur..

- The hard service or host recovery is logged.
- [Event handlers](#) are executed (if you defined any) to handle the hard recovery for the service or host. (Before any event handler is executed, the `$HOSTSTATETYPE$` or `$SERVICESTATETYPE$` macro is set to "**HARD**").
- Contacts will be notified of the service or host recovery (if the [notification logic](#) allows it).

If a hard state change has NOT occurred *and* the service or host is in a non-OK state the following things will occur..

- Contacts will be re-notified of the service or host problem (if the [notification logic](#) allows it).

If a hard state change has NOT occurred *and* the service or host is in an OK state nothing happens. This is because the service or host is in an OK state and was the last time it was checked as well.

---

# Time Periods

or...

"Is This a Good Time?"

---

## Introduction

Time periods allow you to have greater control over when service checks may be run, when host and service notifications may be sent out, and when contacts may receive notifications. With this newly added power come some potential problems, as I will describe later. I was initially very hesitant to introduce time periods because of these snafus. I'll leave it up to you to decide what it right for your particular situation...

## How Time Periods Work With Service Checks

Without the implementation of time periods, Nagios would monitor all services that you had defined 24 hours a day, 7 days a week. While this is fine for most services that need monitoring, it doesn't work out so well for others. For instance, do you really need to monitor printers all the time when they're really only used during normal business hours? Perhaps you have development servers which you would prefer to have up, but aren't "mission critical" and therefore don't have to be monitored for problems over the weekend. Time period definitions now allow you to have more control over when such services may be checked...

The `<check_period>` argument of each service definition allows you to specify a time period that tells Nagios when the service can be checked. When Nagios attempts to reschedule a service check, it will make sure that the next check falls within a valid time range within the defined time period. If it doesn't, Nagios will adjust the next service check time to coincide with the next "valid" time in the specified time period. This means that the service may not get checked again for another hour, day, or week, etc.

## Potential Problems With Service Checks

If you use time periods which do not cover a 24x7 range, you *will* run into problems, especially if a service (or its corresponding host) is down when the check is delayed until the next valid time in the time period. Here are some of those problems...

1. Contacts will not get re-notified of problems with a service until the next service check can be run.
2. If a service recovers during a time that has been excluded from the check period, contacts will not be notified of the recovery.
3. The status of the service will appear unchanged (in the status log and CGI) until it can be checked next.
4. If all services associated with a particular host are on the same check time period, host problems or recoveries will not be recognized until one of the services can be checked (and therefore notifications may be delayed or not get sent out at all).

Limiting the service check period to anything other than a 24 hour a day, 7 days a week basis can cause a lot of problems. Well, not really problems so much as annoyances and inaccuracies... Unless you have good reason to do so, I would *strongly* suggest that you set the `<check_period>` argument of each service definition to a "24x7" type of time period.

## How Time Periods Work With Contact Notifications

Probably the best use of time periods is to control when notifications can be sent out to contacts. By using the `<service_notification_period>` and `<host_notification_period>` arguments in contact definitions, you're able to essentially define an "on call" period for each contact. Note that you can specify different

time periods for host and service notifications. This is helpful if you want host notifications to go out to the contact any day of the week, but only have service notifications get sent to the contact on weekdays. It should be noted that these two notification periods should cover *any time* that the contact can be notified. You can control notification times for specific services and hosts on a one-by-one basis as follows...

By setting the `<notification_period>` argument of the host definition, you can control when Nagios is allowed to send notifications out regarding problems or recoveries for that host. When a host notification is about to get sent out, Nagios will make sure that the current time is within a valid range in the `<notification_period>` time period. If it is a valid time, then Nagios will attempt to notify each contact of the host problem. Some contacts may not receive the host notification if their `<host_notification_period>` does not allow for host notifications at that time. If the time is *not* valid within the `<notification_period>` defined for the host, Nagios will not send the notification out to *any* contacts.

You can control notification times for services in a similar manner to host notification times. By setting the `<notification_period>` argument of the service definition, you can control when Nagios is allowed to send notifications out regarding problems or recoveries for that service. When a service notification is about to get sent out, Nagios will make sure that the current time is within a valid range in the `<notification_period>` time period. If it is a valid time, then Nagios will attempt to notify each contact of the service problem. Some contacts may not receive the service notification if their `<svc_notification_period>` does not allow for service notifications at that time. If the time is *not* valid within the `<notification_period>` defined for the service, Nagios will not send the notification out to *any* contacts.

### **Potential Problems With Contact Notifications**

There aren't really any major problems that you'll run into with using time periods to create custom contact notification times. You do, however, need to be aware that contacts may not always be notified of a service or host problem or recovery. If the time isn't right for both the host or service notification period and the contact notification period, the notification won't go through. Once you weigh the potential problems of time-restricted notifications against your needs, you should be able to come up with a configuration that works well for your situation.

### **Conclusion**

Time periods allow you to have greater control of how Nagios performs its monitoring and notification functions, but can lead to problems. If you are unsure of what type of time periods to implement, or if you are having problems with your current implementation, I would suggest using "24x7" time periods (where all times are valid for each day of the week). Feel free to contact me if you have questions or are running into problems.

---

# Event Handlers

---

## Introduction

Event handlers are optional commands that are executed whenever a host or service state change occurs. An obvious use for event handlers (especially with services) is the ability for Nagios to proactively fix problems before anyone is notified. Another potential use for event handlers might be to log service or host events to an external database.

## Event Handler Types

There are two main types of event handlers than can be defined - service event handlers and host event handlers. Event handler commands are (optionally) defined in each host and service definition. Because these event handlers are only associated with particular services or hosts, I will call these "local" event handlers. If a local event handler has been defined for a service or host, it will be executed when that host or service changes state.

You may also specify global event handlers that should be run for *every* host or service state change by using the [global\\_host\\_event\\_handler](#) and [global\\_service\\_event\\_handler](#) options in your main configuration file. Global event handlers are run immediately *prior* to running a local service or host event handler.

## When Are Event Handler Commands Executed?

Service and host event handler commands are executed when a service or host:

- is in a "soft" error state
- initially goes into a "hard" error state
- recovers from a "soft" or "hard" error state

What are "soft" and "hard" states you ask? They are described [here](#).

## Event Handler Execution Order

Global event handlers are executed before any local event handlers that you have configured for specific hosts or services.

## Writing Event Handler Commands

In most cases, event handler commands will be shell or perl scripts. At a minimum, the scripts should take the following [macros](#) as arguments:

Service event handler macros: `$SERVICESTATE$, $SERVICESTATETYPE$, $SERVICEATTEMPT$`

Host event handler macros: `$HOSTSTATE$, $HOSTSTATETYPE$, $HOSTATTEMPT$`

The scripts should examine the values of the arguments passed in and take any necessary action based upon those values. The best way to understand how event handlers should work is to see an example. Lucky for you, one is provided [below](#). There are also some sample event handler scripts included in the `eventhandlers/` subdirectory of the Nagios distribution. Some of these sample scripts demonstrate the use of [external commands](#) to implement [redundant monitoring hosts](#).

## Permissions For Event Handler Commands

Any event handler commands you configure will execute with the same permissions as the user under which Nagios is running on your machine. This presents a problem with scripts that attempt to restart system services, as root privileges are generally required to do these sorts of tasks.

Ideally you should evaluate the types of event handlers you will be implementing and grant just enough permissions to the Nagios user for executing the necessary system commands. You might want to try using [sudo](#) to accomplish this. Implementation of this is your job, so read the docs and decide if its what you need.

### **Debugging Event Handler Commands**

When you are debugging event handler commands, I would highly recommend that you enable logging of [service retries](#), [host retries](#), and [event handler commands](#). All of these logging options are configured in the [main configuration file](#). Enabling logging for these options will allow you to see exactly when and why event handler commands are being executed.

When you're done debugging your event handler commands you'll probably want to disable logging of service and host retries. They can fill up your log file fast, but if you have enabled [log rotation](#) you might not care.

### **Service Event Handler Example**

The example below assumes that you are monitoring the HTTP server on the local machine and have specified **restart-httpd** as the event handler command for the HTTP service definition. Also, I will be assuming that you have set the `<max_check_attempts>` option for the service to be a value of 4 or greater (i.e. the service is checked 4 times before it is considered to have a real problem). An example service definition (w/ only the fields we discuss) might look like this...

```
define service{
    host_name             somehost
    service_description   HTTP
    max_check_attempts    4
    event_handler         restart-httpd
    ...other service variables...
}
```

Once the service has been defined with an event handler, we must define that event handler as a command. Notice the macros in the command line that I am passing to the event handler - these are important!

```
define command{
    command_name    restart-httpd
    command_line    /usr/local/nagios/libexec/eventhandlers/restart-httpd $SERVICESTATES $SERVICESTATETYPES $SERVICEATTEMPT$
}
```

Now, let's actually write the event handler script (this is the `/usr/local/nagios/libexec/eventhandlers/restart-httpd` file).

```
#!/bin/sh
#
# Event handler script for restarting the web server on the local machine
#
# Note: This script will only restart the web server if the service is
#       retried 3 times (in a "soft" state) or if the web service somehow
#       manages to fall into a "hard" error state.
#

# What state is the HTTP service in?
case "$1" in
OK)
# The service just came back up, so don't do anything...
;;
WARNING)
# We don't really care about warning states, since the service is probably still running...
;;
UNKNOWN)
# We don't know what might be causing an unknown error, so don't do anything...
;;
CRITICAL)
# Aha! The HTTP service appears to have a problem - perhaps we should restart the server...

# Is this a "soft" or a "hard" state?
case "$2" in

# We're in a "soft" state, meaning that Nagios is in the middle of retrying the
# check before it turns into a "hard" state and contacts get notified...
SOFT)

# What check attempt are we on? We don't want to restart the web server on the first
# check, because it may just be a fluke!
case "$3" in

# Wait until the check has been tried 3 times before restarting the web server.
# If the check fails on the 4th time (after we restart the web server), the state
# type will turn to "hard" and contacts will be notified of the problem.
# Hopefully this will restart the web server successfully, so the 4th check will
# result in a "soft" recovery. If that happens no one gets notified because we
# fixed the problem!
3)
echo -n "Restarting HTTP service (3rd soft critical state)..."
# Call the init script to restart the HTTPD server
/etc/rc.d/init.d/httpd restart
;;
esac
;;

# The HTTP service somehow managed to turn into a hard error without getting fixed.
# It should have been restarted by the code above, but for some reason it didn't.
# Let's give it one last try, shall we?
# Note: Contacts have already been notified of a problem with the service at this
# point (unless you disabled notifications for this service)
HARD)
echo -n "Restarting HTTP service..."
# Call the init script to restart the HTTPD server
/etc/rc.d/init.d/httpd restart
;;
esac
;;
esac
exit 0
```

The sample script provided above will attempt to restart the web server on the local machine in two different instances - after the HTTP service is being retried for the 3rd time (in an "soft" error state) and after the service falls into a "hard" state. The "hard" state situation shouldn't really occur, since the script should restart the service when its still in a "soft" state (i.e. the 3rd check retry), but its left as a fallback anyway.

It should be noted that the service event handler will only be execute the first time that the service falls into a "hard" state. This will prevent Nagios from continuously executing the script to restart the web server when it is in a "hard" state.





# External Commands

---

## Introduction

Nagios can process commands from external applications (including CGIs - see the [command CGI](#) for an example) and alter various aspects of its monitoring functions based on the commands it receives.

## Enabling External Commands

By default, Nagios *does not* check for or process any external commands. If you want to enable external command processing, you'll have to do the following...

- Enable external command checking with the [check\\_external\\_commands](#) option
- Set the frequency of command checks with the [command\\_check\\_interval](#) option
- Specify the location of the command file with the [command\\_file](#) option. Its best to put the external command file in its own directory (i.e. `/usr/local/nagios/var/rw`).
- Setup proper permissions on the directory containing the external command file. Details on how to do this can be found [here](#).

## When Does Nagios Check For External Commands?

- At regular intervals specified by the [command\\_check\\_interval](#) option in the main configuration file
- Immediately after [event handlers](#) are executed. This is in addition to the regular cycle of external command checks and is done to provide immediate action if an event handler submits commands to Nagios.

## Using External Commands

External commands can be used to accomplish a variety of things while Nagios is running. Example of what can be done include temporarily disabling notifications for services and hosts, temporarily disabling service checks, forcing immediate service checks, adding comments to hosts and services, etc.

## Command Format

External commands that are written to the [command file](#) have the following format...

***[time] command\_id;command\_arguments***

...where *time* is the time (in *time\_t* format) that the external application or CGI committed the external command to the command file. Some of the commands that are available are described in the table below, along with their *command\_id* and a description of their *command\_arguments*.

## Implemented Commands

A full listing of external commands that can be used (along with examples of how to use them) can be found online at the following URL:

<http://www.nagios.org/developerinfo/externalcommands/>

---

# Indirect Host and Service Checks

---

## Introduction

Chances are, many of the services that you're going to be monitoring on your network can be checked directly by using a plugin on the host that runs Nagios. Examples of services that can be checked directly include availability of web, email, and FTP servers. These services can be checked directly by a plugin from the Nagios host because they are publicly accessible resources. However, there are a number of things you may be interested in monitoring that are not as publicly accessible as other services. These "private" resources/services include things like disk usage, processor load, etc. on remote machines. Private resources like these cannot be checked without the use of an intermediary agent. Service checks which require an intermediary agent of some kind to actually perform the check are called *indirect* checks.

Indirect checks are useful for:

- Monitoring "local" resources (such as disk usage, processor load, etc.) on remote hosts
- Monitoring services and hosts behind firewalls
- Obtaining more realistic results from checks of time-sensitive services between remote hosts (i.e. ping response times between two remote hosts)

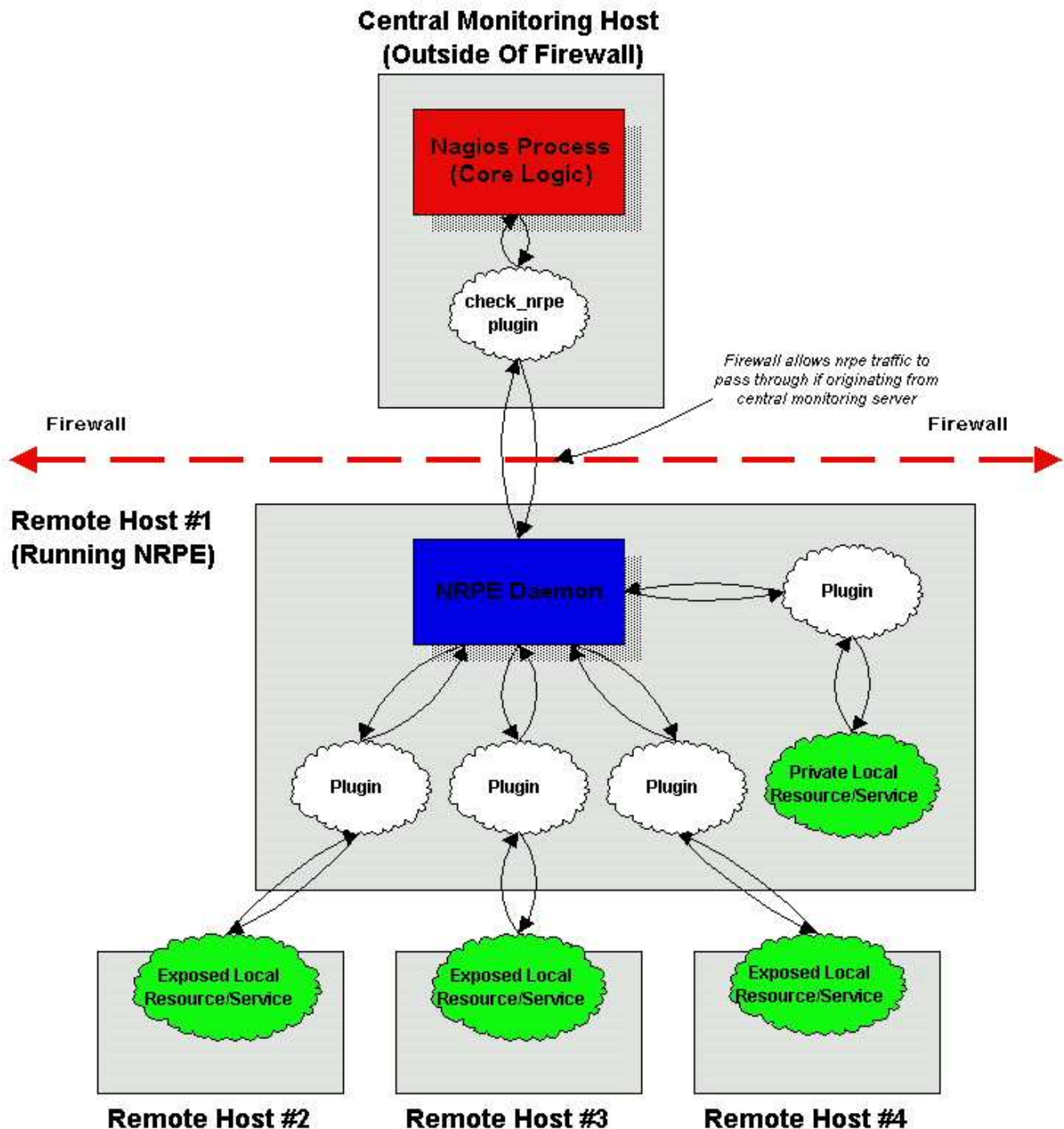
There are several methods for performing indirect active checks ([passive checks](#) are not discussed here), but I will only talk about how they can be done by using the [nrpe](#) addon.

## Indirect Service Checks

The diagram below shows how indirect service checks work. Click the image for a larger version...

# Indirect Service Checks

Last Updated: 07-12-2001



## Multiple Indirected Service Checks

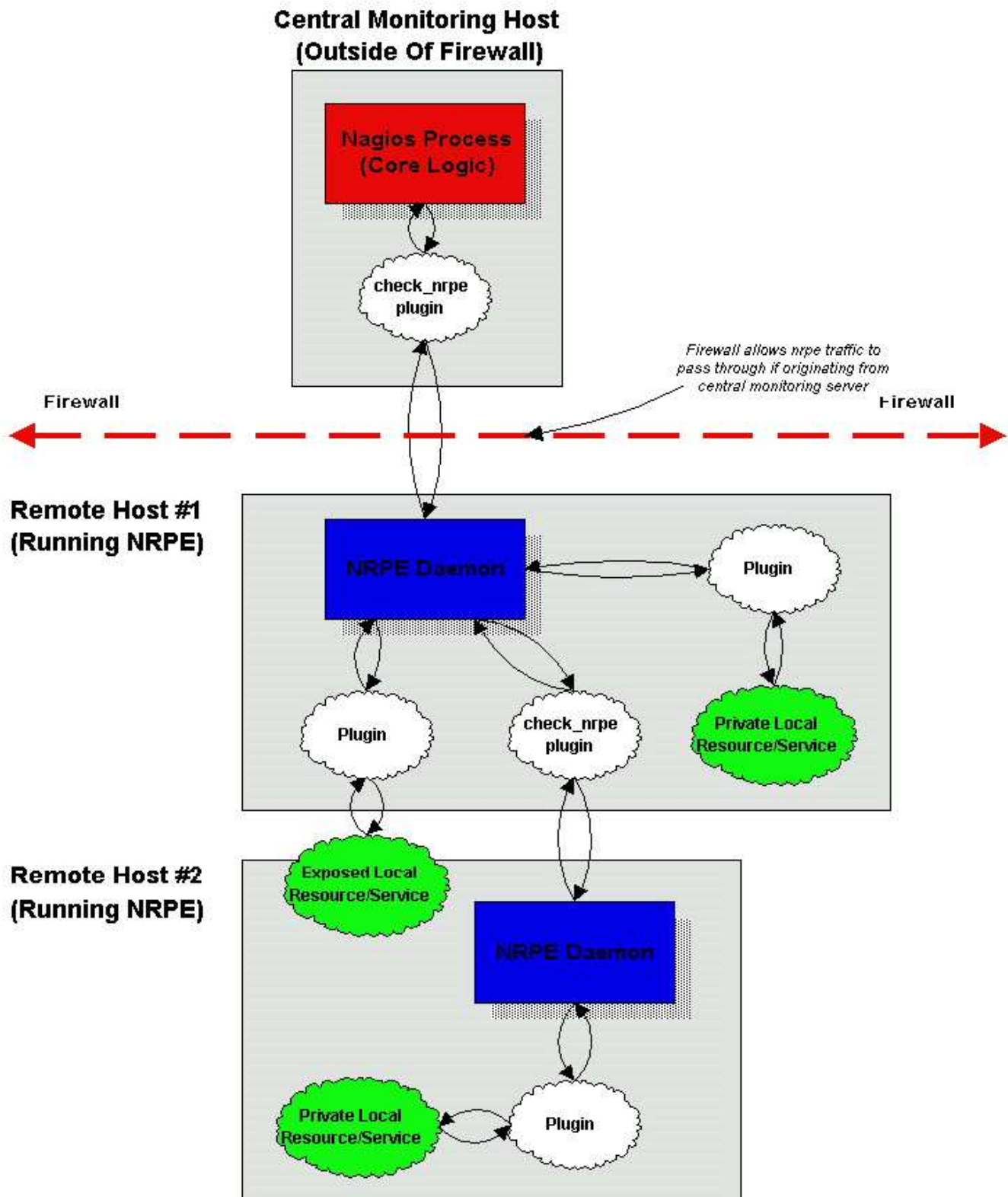
If you are monitoring servers that lie behind a firewall (and the host running Nagios is outside that firewall), checking services on those machines can prove to be a bit of a pain. Chances are that you are blocking most incoming traffic that would normally be required to perform the monitoring. One solution for performing active checks ([passive checks](#) could also be used) on the hosts behind the firewall would be to poke a tiny hold in the firewall filters that allow the Nagios host to make calls to the

*nrpe* daemon on one host inside the firewall. The host inside the firewall could then be used as an intermediary in performing checks on the other servers inside the firewall.

The diagram below show how multiple indirect service checks work. Notice how the *nrpe* daemon is running on hosts #1 and #2. The copy that runs on host #2 is used to allow the *nrpe* agent on host #1 to perform a check of a "private" service on host #2. "Private" services are things like process load, disk usage, etc. that are not directly exposed like SMTP, FTP, and web services. Click on the diagram for a larger image...

# Multiple Indirected Service Checks

Last Updated: 07-21-2001



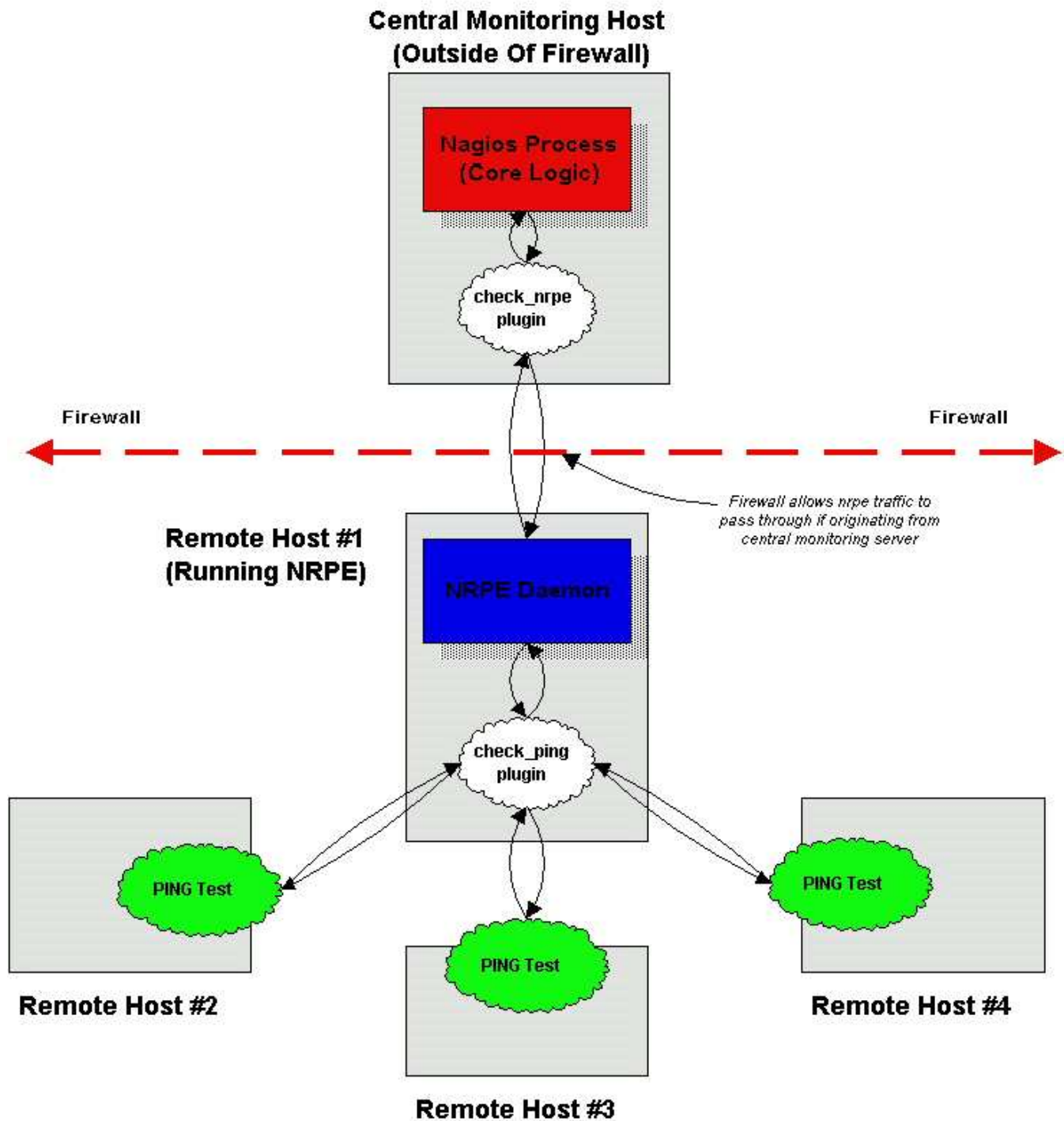
## Indirect Host Checks

Indirect host checks work on the same principle as indirect service checks. Basically, the plugin used in the host check command asks an intermediary agent (i.e. a daemon running on a remote host) to perform the host check for it. Indirect host checks are useful when the remote hosts being monitored are located behind a firewall and you want to restrict inbound monitoring traffic to a particular machine. That machine (remote host #1 in the diagram below) performs will perform the host check and return the results back to the top level *check\_nrpe* plugin (on the central server). It should be noted that with this setup comes potential problems. If remote host #1 goes down, the *check\_nrpe* plugin will not be able to contact the *nrpe* daemon and Nagios will believe that remote hosts #2, #3, and #4 are down, even though this may not be the case. If host #1 is your firewall machine, then the problem isn't really an issue because Nagios will detect that it is down and mark hosts #2, #3, and #4 as being unreachable.

The diagram below shows how an indirect host check can be performed by using the *nrpe* daemon and *check\_nrpe* plugin. Click the image for a larger version.

# Indirect Host Checks

Last Updated: 07-12-2001





# Passive Host and Service Checks

---

## Introduction

One of the features of Nagios is that it can process host and service check results that are submitted by external applications. Host and service checks which are performed and submitted to Nagios by external apps are called *passive* checks. Passive checks can be contrasted with *active* checks, which are host or service checks that have been initiated by Nagios.

## Why The Need For Passive Checks?

Passive checks are useful for monitoring services that are:

- located behind a firewall, and can therefore not be checked actively from the host running Nagios
- asynchronous in nature and can therefore not be actively checked in a reliable manner (e.g. SNMP traps, security alerts, etc.)

Passive host and service checks are also useful when configured a [distributed monitoring setup](#).

## Passive Service Checks vs. Passive Host Checks

Passive host and service checks function in a similar manner, but there are some important limitations in regards to passive host checks. Read [below](#) for more information about the limitations with passive host checks.

## How Do Passive Service Checks Work?

The only real difference between active and passive checks is that active checks are initiated by Nagios, while passive checks are performed by external applications. Once an external application has performed a service check (either actively or by having received an synchronous event like an SNMP trap or security alert), it submits the results of the service "check" to Nagios through the [external command file](#).

The next time Nagios processes the contents of the external command file, it will place the results of all passive service checks into a queue for later processing. The same queue that is used for storing results from active checks is also used to store the results from passive checks.

Nagios will periodically execute a [service reaper event](#) and scan the service check result queue. Each service check result, regardless of whether the check was active or passive, is processed in the same manner. The service check logic is exactly the same for both types of checks. This provides a seamless method for handling both active and passive service check results.

## How Do External Apps Submit Service Check Results?

External applications can submit service check results to Nagios by writing a `PROCESS_SERVICE_CHECK_RESULT` [external command](#) to the [external command file](#).

The format of the command is as follows:

```
[<timestamp>]  
PROCESS_SERVICE_CHECK_RESULT;<host_name>;<description>;<return_code>;<plugin_output>
```

where...

- *timestamp* is the time in time\_t format (seconds since the UNIX epoch) that the service check was performed (or submitted). Please note the single space after the right bracket.
- *host\_name* is the short name of the host associated with the service in the service definition
- *description* is the description of the service as specified in the service definition
- *return\_code* is the return code of the check (0=OK, 1=WARNING, 2=CRITICAL, 3=UNKNOWN)
- *plugin\_output* is the text output of the service check (i.e. the plugin output)

Note that in order to submit service checks to Nagios, a service must have already been defined in the [object configuration file](#)! Nagios will ignore all check results for services that had not been configured before it was last (re)started.

If you only want passive results to be provided for a specific service (i.e. active checks should not be performed), simply set the *active\_checks\_enabled* member of the service definition to 0. This will prevent Nagios from ever actively performing a check of the service. Make sure that the *passive\_checks\_enabled* member of the service definition is set to 1. If it isn't, Nagios won't process passive checks for the service!

An example shell script of how to submit passive service check results to Nagios can be found in the documentation on [volatile services](#).

### **Submitting Passive Service Check Results From Remote Hosts**

If an application that resides on the same host as Nagios is sending passive service check results, it can simply write the results directly to the external command file as outlined above. However, applications on remote hosts can't do this so easily. In order to allow remote hosts to send passive service check results to the host that runs Nagios, I've developed the [nsca](#) addon. The addon consists of a daemon that runs on the Nagios hosts and a client that is executed from remote hosts. The daemon will listen for connections from remote clients, perform some basic validation on the results being submitted, and then write the check results directly into the external command file (as described above). More information on the nsca addon can be found [here](#)...

### **Using Both Active And Passive Service Checks**

Unless you're implementing a [distributed monitoring](#) environment with the central server accepting only passive service checks (and not performing any active checks), you'll probably be using both types of checks in your setup. As mentioned before, active checks are more suited for services that lend themselves to periodic checks (availability of an FTP or web server, etc), whereas passive checks are better off at handling asynchronous events that occur at variable intervals (security alerts, etc.).

The image below gives a visual representation of how active and passive service checks can both be used to monitor network resources (click on the image for a larger version).

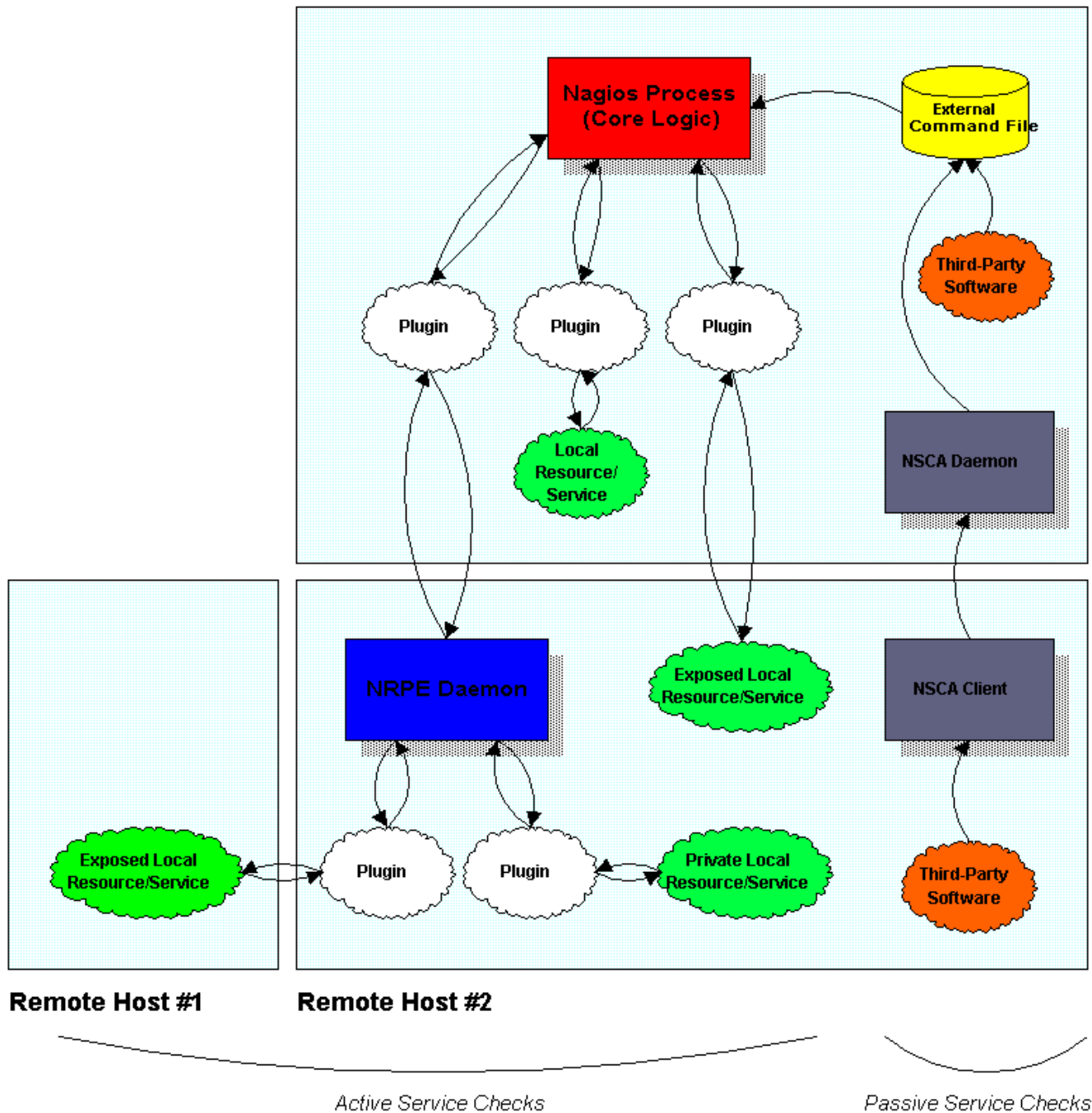
The orange bubbles on the right side of the image are third-party applications that submit passive check results to Nagios' external command file. One of the applications resides on the same host as Nagios, so it can write directly to the command file. The other application resides on a remote host and makes use of the nsca client program and daemon to transfer the passive check results to Nagios.

The items on the left side of the image represent active service checks that Nagios is performing. I've shown how the checks can be made for local resources (disk usage, etc.), "exposed" resources on remote hosts (web server, FTP server, etc.), and "private" resources on remote hosts (remote host disk usage, processor load, etc.). In this example, the private resources on the remote hosts are actually checked by making use of the [nrpe](#) addon, which facilitates the execution of plugins on remote hosts.

# Using Active And Passive Checks Together

Last Updated: 07-21-2001

## Monitoring Host



## How Do Passive Host Checks Work?

Passive host checks work in a similar manner to passive service checks. Once an external application has performed a host check, it submits the results of that host "check" to Nagios through the [external command file](#). The next time Nagios processes the contents of the external command file, it will process the host check result that was submitted.

**WARNING!** Passive host checks have some limitations. Unlike active host checks, Nagios does not attempt to determine whether or host is DOWN or UNREACHABLE with passive checks. Rather, Nagios takes the passive check result to be the actual state the host is in and doesn't try to determine the actual state. In contrast, Nagios attempts to determine the proper status (DOWN or UNREACHABLE) for hosts that are not UP when the host check is active (initiated by Nagios). This can cause problems if you are submitting passive checks from a remote host or you have a [distributed monitoring setup](#) where the parent/child host relationships are different. See the documentation on [host reachability](#) for more information on how DOWN and UNREACHABLE states are determined for active host checks.

### How Do External Apps Submit Host Check Results?

External applications can submit host check results to Nagios by writing a `PROCESS_HOST_CHECK_RESULT` [external command](#) to the [external command file](#).

The format of the command is as follows:

**[<timestamp>] PROCESS\_HOST\_CHECK\_RESULT;<host\_name>;<host\_status>;<plugin\_output>**

where...

- *timestamp* is the time in time\_t format (seconds since the UNIX epoch) that the host check was performed (or submitted). Please note the single space after the right bracket.
- *host\_name* is the short name of the host (as defined in the host definition)
- *host\_status* is the status of the host (0=UP, 1=DOWN, 2=UNREACHABLE)
- *plugin\_output* is the text output of the host check

Note that in order to submit host checks to Nagios, a host must have already been defined in the [object configuration file](#)! Nagios will ignore all check results for hosts that had not been configured before it was last (re)started.

### Submitting Passive Host Check Results From Remote Hosts

If an application that resides on the same host as Nagios is sending passive service check results, it can simply write the results directly to the external command file as outlined above. However, applications on remote hosts can't do this so easily. In order to allow remote hosts to send passive host check results to the host that runs Nagios, you can use the [nsca](#) addon. The addon consists of a daemon that runs on the Nagios hosts and a client that is executed from remote hosts. The daemon will listen for connections from remote clients, perform some basic validation on the results being submitted, and then write the check results directly into the external command file (as described above). More information on the nsca addon can be found [here](#).

---

# Volatile Services

---

## Introduction

Nagios has the ability to distinguish between "normal" services and "volatile" services. The *is\_volatile* option in each service definition allows you to specify whether a specific service is volatile or not. For most people, the majority of all monitored services will be non-volatile (i.e. "normal"). However, volatile services can be very useful when used properly...

## What Are They Useful For?

Volatile services are useful for monitoring...

- things that automatically reset themselves to an "OK" state each time they are checked
- events such as security alerts which require attention every time there is a problem (and not just the first time)

## What's So Special About Volatile Services?

Volatile services differ from "normal" services in three important ways. *Each time* they are checked when they are in a **hard** non-OK state, and the check returns a non-OK state (i.e. no state change has occurred)...

- the non-OK service state is logged
- contacts are notified about the problem (if that's **what should be done**)
- the **event handler** for the service is run (if one has been defined)

These events normally only occur for services when they are in a non-OK state and a hard state change has just occurred. In other words, they only happen the first time that a service goes into a non-OK state. If future checks of the service result in the same non-OK state, no hard state change occurs and none of the events mentioned take place again.

## The Power Of Two

If you combine the features of volatile services and **passive service checks**, you can do some very useful things. Examples of this include handling SNMP traps, security alerts, etc.

How about an example... Let's say you're running **PortSentry** to detect port scans on your machine and automatically firewall potential intruders. If you want to let Nagios know about port scans, you could do the following..

### **In Nagios:**

- Configure a service called *Port Scans* and associate it with the host that PortSentry is running on.
- Set the *max\_check\_attempts* option in the service definition to 1. This will tell Nagios to immediately force the service into a **hard state** when a non-OK state is reported.
- Either set the *active\_checks\_enabled* option to 0 or set the *check\_time* option in the service definition to a **timeperiod** that contains *no* valid time ranges. Doing either of these will prevent Nagios from ever actively checking the service. Even though the service check will get scheduled, it will never actually be checked.

### **In PortSentry:**

- Edit your PortSentry configuration file (portsentry.conf), define a command for the **KILL\_RUN\_CMD** directive as follows:  
**KILL\_RUN\_CMD="/usr/local/Nagios/libexec/eventhandlers/submit\_check\_result <host\_name> 'Port Scans' 2 'Port scan from host \$TARGET\$ on port \$PORT\$. Host has been firewalled.'"**  
 Make sure to replace <host\_name> with the short name of the host that the service is associated with.

Create a shell script in the `/usr/local/nagios/libexec/eventhandlers` directory named `submit_check_result`. The contents of the shell script should be something similar to the following...

```
#!/bin/sh

# Write a command to the Nagios command file to cause
# it to process a service check result

echocmd="/bin/echo"

CommandFile="/usr/local/nagios/var/rw/nagios.cmd"

# get the current date/time in seconds since UNIX epoch
datetime='date +%s'

# create the command line to add to the command file
cmdline="[$datetime] PROCESS_SERVICE_CHECK_RESULT;$1;$2;$3;$4"

# append the command to the end of the command file
'$echocmd $cmdline >> $CommandFile'
```

Note that if you are running PortSentry as root, you will have to make additions to the script to reset file ownership and permissions so that Nagios and the CGIs can read/modify the command file.

So what happens when PortSentry detects a port scan on the machine?

- It blocks the host (this is a function of the PortSentry software)
  - It executes the `submit_check_result` shell script to send the security alert info to Nagios
  - Nagios reads the command file, recognized the port scan entry as a passive service check
  - Nagios processes the results of the service by logging the CRITICAL state, sending notifications to contacts (if configured to do so), and executes the event handler for the *Port Scans* service (if one is defined)
-

# Service and Host Result Freshness Checks

---

## Introduction

Nagios supports a feature that does "freshness" checking on the results of host and service checks. This feature is useful when you want to ensure that [passive checks](#) are being received as frequently as you want. Although freshness checking can be used in a number of situations, it is primarily useful when attempting to configure a [distributed monitoring environment](#).

The purpose of "freshness" checking is to ensure that host and service checks are being provided passively by external applications on a regular basis. If the results of a particular host or service check (for which freshness checking has been enabled) is determined to be "stale", Nagios will force an active check of that host or service.

## Host vs. Service Freshness Checking

The documentation below discusses service freshness checking. Host freshness checking (which is not documented separately) works in a similar way to service freshness checking - except, of course, that its for hosts instead of services. If you need to configure host freshness checking, adjust the directions given below appropriately.

## Configuring Service Freshness Checking

Before you configure per-service freshness threshold, you must enable freshness checking using the [check\\_service\\_freshness](#) and [service\\_freshness\\_check\\_interval](#) directives in the main config file. If you were configuring host freshness checking, you would use the [check\\_host\\_freshness](#) and [host\\_freshness\\_check\\_interval](#) directives.

So how do you go about enabling freshness checking for a particular service? You need to configure [service definitions](#) as follows.

- The **check\_freshness** option in the service definition should be set to 1. This enables "freshness" checking for the service.
- The **freshness\_threshold** option in the service definition should be set to a value (in seconds) which reflects how "fresh" the results for the service should be.
- The **check\_command** option in the service definition should reflect valid command that should be used to actively check the service when it is detected as being "stale".
- The **normal\_check\_interval** option in the service definition needs to be greater than zero (0) if the **freshness\_threshold** option is setup to zero (0).
- The **check\_period** option in the service definitions needs to be set to a valid timeperiod. The times allowed by the specified timeperiod determine when freshness checks can be performed for the service.

## How The Freshness Threshold Works

Nagios periodically checks the "freshness" of the results for all services that have freshness checking enabled. The *freshness\_threshold* option in each service definition is used to determine how "fresh" the results for each service should be. For example, if you set the *freshness\_threshold* option to 60 for one of your services, Nagios will consider that service to be "stale" if its results are older than 60 seconds (1 minute). If you do not specify a value for the *freshness\_threshold* option (or you set it to zero), Nagios will automatically calculate a "freshness" threshold to use by looking at either the *normal\_check\_interval* or *retry\_check\_interval* options (depending on what [type of state](#) the service is currently in).



## What Happens When A Service Check Result Becomes "Stale"

If the check results of a service are found to be "stale" (as described above), Nagios will force an active check of the service by executing the command specified by the *check\_command* option in the service definition. It is important to note that an active service check which is being forced because the service was detected as being "stale" gets executed *even if active service checks are disabled on a program-wide or service-specific basis*.

## Working With Passive-Only Checks

As I mentioned earlier, freshness checking is of most use when you are dealing with services that get their results from [passive checks](#). More often than not (as in the case with [distributed monitoring setups](#)), these services may not be getting *all* of their results from passive checks - no results are obtained from active checks.

An example of a passive-only service might be one that reports the status of your nightly backup jobs. Perhaps you have an external script that submit the results of the backup job to Nagios once the backup is completed. In this case, all of the checks/results for the service are provided by an external application using passive checks. In order to ensure that the status of the backup job gets reported every day, you may want to enable freshness checking for the service. If the external script doesn't submit the results of the backup job, you can have Nagios fake a critical result by doing something like this...

Here's what the definition for the service might look like (some required options are omitted)...

```
define service{
    host_name          backup-server
    service_description ArcServe Backup Job
    active_checks_enabled 0           ; active checks are NOT enabled
    passive_checks_enabled 1          ; passive checks are enabled (this is how results are reported)
    check_freshness     1
    freshness_threshold 93600         ; 26 hour threshold, since backups may not always finish at the same time
    check_command        no-backup-report ; this command is run only if the service results are "stale"
    ...other options...
}
```

Notice that active checks are disabled for the service. This is because the results for the service are only made by an external application using passive checks. Freshness checking is enabled and the freshness threshold has been set to 26 hours. This is a bit longer than 24 hours because backup jobs sometimes run late from day to day (depending on how much data there is to backup, how much network traffic is present, etc.). The *no-backup-report* command is executed only if the results of the service are determined to be "stale". The definition of the *no-backup-report* command might look like this...

```
define command{
    command_name    no-backup-report
    command_line    /usr/local/nagios/libexec/nobackupreport.sh
}
```

The **nobackupreport.sh** script in your */usr/local/nagios/libexec* directory might look something like this:

```
#!/bin/sh

/bin/echo "CRITICAL: Results of backup job were not reported!"

exit 2
```

If Nagios detects that the service results are stale, it will run the **no-backup-report** command as an active service check (even though active checks are disabled for this specific service - remember that this is a special case). This causes the */usr/local/nagios/libexec/nobackupreport.sh* script to be executed, which returns a critical state. The service goes into a critical state (if it isn't already there) and someone will probably get notified of the problem.





# Distributed Monitoring

---

## Introduction

Nagios can be configured to support distributed monitoring of network services and resources. I'll try to briefly explain how this can be accomplished...

## Goals

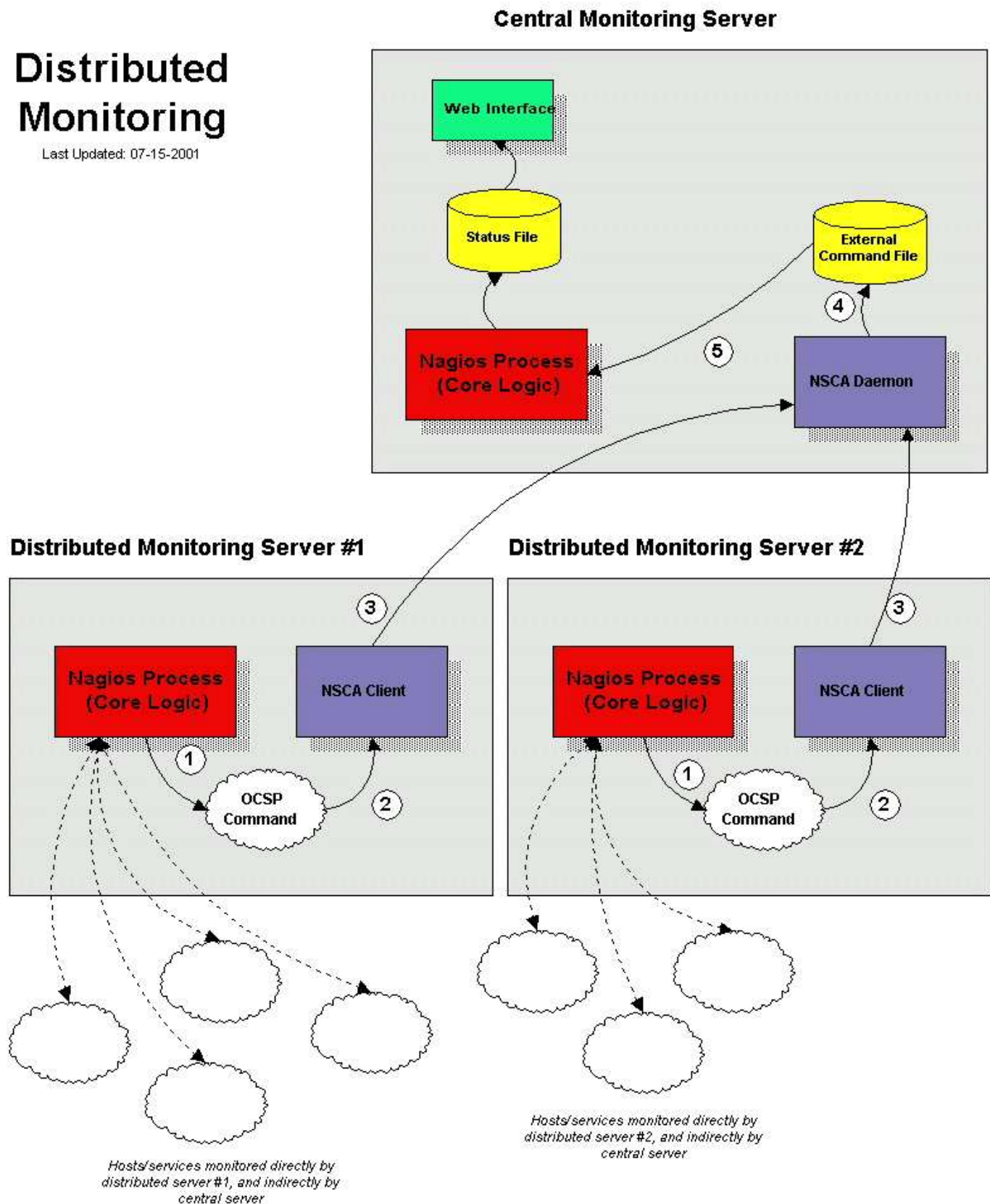
The goal in the distributed monitoring environment that I will describe is to offload the overhead (CPU usage, etc.) of performing service checks from a "central" server onto one or more "distributed" servers. Most small to medium sized shops will not have a real need for setting up such an environment. However, when you want to start monitoring hundreds or even thousands of *hosts* (and several times that many services) using Nagios, this becomes quite important.

## Reference Diagram

The diagram below should help give you a general idea of how distributed monitoring works with Nagios. I'll be referring to the items shown in the diagram as I explain things...

# Distributed Monitoring

Last Updated: 07-15-2001



## Central Server vs. Distributed Servers

When setting up a distributed monitoring environment with Nagios, there are differences in the way the central and distributed servers are configured. I'll show you how to configure both types of servers and explain what effects the changes being made have on the overall monitoring. For starters, let's describe

the purpose of the different types of servers...

The function of a *distributed server* is to actively perform checks all the services you define for a "cluster" of hosts. I use the term "cluster" loosely - it basically just mean an arbitrary group of hosts on your network. Depending on your network layout, you may have several clusters at one physical location, or each cluster may be separated by a WAN, its own firewall, etc. The important thing to remember to that for each cluster of hosts (however you define that), there is one distributed server that runs Nagios and monitors the services on the hosts in the cluster. A distributed server is usually a bare-bones installation of Nagios. It doesn't have to have the web interface installed, send out notifications, run event handler scripts, or do anything other than execute service checks if you don't want it to. More detailed information on configuring a distributed server comes later...

The purpose of the *central server* is to simply listen for service check results from one or more distributed servers. Even though services are occasionally actively checked from the central server, the active checks are only performed in dire circumstances, so let's just say that the central server only accepts passive check for now. Since the central server is obtaining [passive service check](#) results from one or more distributed servers, it serves as the focal point for all monitoring logic (i.e. it sends out notifications, runs event handler scripts, determines host states, has the web interface installed, etc).

### **Obtaining Service Check Information From Distributed Monitors**

Okay, before we go jumping into configuration detail we need to know how to send the service check results from the distributed servers to the central server. I've already discussed how to submit passive check results to Nagios from same host that Nagios is running on (as described in the documentation on [passive checks](#)), but I haven't given any info on how to submit passive check results from other hosts.

In order to facilitate the submission of passive check results to a remote host, I've written the [nsca addon](#). The addon consists of two pieces. The first is a client program (`send_nsca`) which is run from a remote host and is used to send the service check results to another server. The second piece is the `nsca` daemon (`nsca`) which either runs as a standalone daemon or under `inetd` and listens for connections from client programs. Upon receiving service check information from a client, the daemon will submit the check information to Nagios (on the central server) by inserting a `PROCESS_SVC_CHECK_RESULT` command into the [external command file](#), along with the check results. The next time Nagios checks for [external commands](#), it will find the passive service check information that was sent from the distributed server and process it. Easy, huh?

### **Distributed Server Configuration**

So how exactly is Nagios configured on a distributed server? Basically, it's just a bare-bones installation. You don't need to install the web interface or have notifications sent out from the server, as this will all be handled by the central server.

Key configuration changes:

- Only those services and hosts which are being monitored directly by the distributed server are defined in the [object configuration file](#).
- The distributed server has its [enable\\_notifications](#) directive set to 0. This will prevent any notifications from being sent out by the server.
- The distributed server is configured to [obsess over services](#).
- The distributed server has an [ocsp command](#) defined (as described below).

In order to make everything come together and work properly, we want the distributed server to report the results of *all* service checks to Nagios. We could use [event handlers](#) to report *changes* in the state of a service, but that just doesn't cut it. In order to force the distributed server to report all service check results, you must enable the [obsess\\_over\\_services](#) option in the main configuration file and provide a

ocsp\_command to be run after every service check. We will use the ocsp command to send the results of all service checks to the central server, making use of the send\_nsca client and nsca daemon (as described above) to handle the transmission.

In order to accomplish this, you'll need to define an ocsp command like this:

### ocsp\_command=submit\_check\_result

The command definition for the *submit\_check\_result* command looks something like this:

```
define command{
    command_name    submit_check_result
    command_line    /usr/local/nagios/libexec/eventhandlers/submit_check_result $HOSTNAME$ '$SERVICEDESC$' $SERVICESTATE$ '$SERVICEOUTPUT$'
}
```

The *submit\_check\_result* shell script looks something like this (replace *central\_server* with the IP address of the central server):

```
#!/bin/sh

# Arguments:
# $1 = host_name (Short name of host that the service is
# associated with)
# $2 = svc_description (Description of the service)
# $3 = state_string (A string representing the status of
# the given service - "OK", "WARNING", "CRITICAL"
# or "UNKNOWN")
# $4 = plugin_output (A text string that should be used
# as the plugin output for the service checks)
#

# Convert the state string to the corresponding return code
return_code=-1

case "$3" in
    OK)
        return_code=0
        ;;
    WARNING)
        return_code=1
        ;;
    CRITICAL)
        return_code=2
        ;;
    UNKNOWN)
        return_code=-1
        ;;
esac

# pipe the service check info into the send_nsca program, which
# in turn transmits the data to the nsca daemon on the central
# monitoring server

/bin/printf "%s\t%s\t%s\t%s\n" "$1" "$2" "$return_code" "$4" | /usr/local/nagios/bin/send_nsca central_server -c /usr/local/nagios/etc/send_nsca.cfg
```

The script above assumes that you have the send\_nsca program and its configuration file (send\_nsca.cfg) located in the */usr/local/nagios/bin/* and */usr/local/nagios/etc/* directories, respectively.

That's it! We've successfully configured a remote host running Nagios to act as a distributed monitoring server. Let's go over exactly what happens with the distributed server and how it sends service check results to Nagios (the steps outlined below correspond to the numbers in the reference diagram above):

1. After the distributed server finishes executing a service check, it executes the command you defined by the **ocsp\_command** variable. In our example, this is the */usr/local/nagios/libexec/eventhandlers/submit\_check\_result* script. Note that the definition for the *submit\_check\_result* command passed four pieces of information to the script: the name of the host the service is associated with, the service description, the return code from the service check, and the plugin output from the service check.
2. The *submit\_check\_result* script pipes the service check information (host name, description, return code, and output) to the *send\_nsca* client program.
3. The *send\_nsca* program transmits the service check information to the *nsca* daemon on the central monitoring server.
4. The *nsca* daemon on the central server takes the service check information and writes it to the external command file for later pickup by Nagios.
5. The Nagios process on the central server reads the external command file and processes the passive

service check information that originated from the distributed monitoring server.

## **Central Server Configuration**

We've looked at how distributed monitoring servers should be configured, so let's turn to the central server. For all intensive purposes, the central is configured as you would normally configure a standalone server. It is setup as follows:

- The central server has the web interface installed (optional, but recommended)
- The central server has its [enable\\_notifications](#) directive set to 1. This will enable notifications. (optional, but recommended)
- The central server has [active service checks](#) disabled (optional, but recommended - see notes below)
- The central server has [external command checks](#) enabled (required)
- The central server has [passive service checks](#) enabled (required)

There are three other very important things that you need to keep in mind when configuring the central server:

- The central server must have service definitions for *all services* that are being monitored by all the distributed servers. Nagios will ignore passive check results if they do not correspond to a service that has been defined.
- If you're only using the central server to process services whose results are going to be provided by distributed hosts, you can simply disable all active service checks on a program-wide basis by setting the [execute\\_service\\_checks](#) directive to 0. If you're using the central server to actively monitor a few services on its own (without the aid of distributed servers), the [enable\\_active\\_checks](#) option of the definitions for service being monitored by distributed servers should be set to 0. This will prevent Nagios from actively checking those services.

It is important that you either disable all service checks on a program-wide basis or disable the [enable\\_active\\_checks](#) option in the definitions for each service that is monitored by a distributed server. This will ensure that active service checks are never executed under normal circumstances. The services will keep getting rescheduled at their normal check intervals (3 minutes, 5 minutes, etc...), but they won't actually be executed. This rescheduling loop will just continue all the while Nagios is running. I'll explain why this is done in a bit...

That's it! Easy, huh?

## **Problems With Passive Checks**

For all intensive purposes we can say that the central server is relying solely on passive checks for monitoring. The main problem with relying completely on passive checks for monitoring is the fact that Nagios must rely on something else to provide the monitoring data. What if the remote host that is sending in passive check results goes down or becomes unreachable? If Nagios isn't actively checking the services on the host, how will it know that there is a problem?

Fortunately, there is a way we can handle these types of problems...

## **Freshness Checking**

Nagios supports a feature that does "freshness" checking on the results of service checks. More information freshness checking can be found [here](#). This feature gives some protection against situations where remote hosts may stop sending passive service checks into the central monitoring server. The purpose of "freshness" checking is to ensure that service checks are either being provided passively by distributed servers on a regular basis or performed actively by the central server if the need arises. If the service check results provided by the distributed servers get "stale", Nagios can be configured to force

active checks of the service from the central monitoring host.

So how do you do this? On the central monitoring server you need to configure services that are being monitoring by distributed servers as follows...

- The *check\_freshness* option in the service definitions should be set to 1. This enables "freshness" checking for the services.
- The *freshness\_threshold* option in the service definitions should be set to a value (in seconds) which reflects how "fresh" the results for the services (provided by the distributed servers) should be.
- The *check\_command* option in the service definitions should reflect valid commands that can be used to actively check the service from the central monitoring server.

Nagios periodically checks the "freshness" of the results for all services that have freshness checking enabled. The *freshness\_threshold* option in each service definition is used to determine how "fresh" the results for each service should be. For example, if you set this value to 300 for one of your services, Nagios will consider the service results to be "stale" if they're older than 5 minutes (300 seconds). If you do not specify a value for the *freshness\_threshold* option, Nagios will automatically calculate a "freshness" threshold by looking at either the *normal\_check\_interval* or *retry\_check\_interval* options (depending on what [type of state](#) the service is in). If the service results are found to be "stale", Nagios will run the service check command specified by the *check\_command* option in the service definition, thereby actively checking the service.

Remember that you have to specify a *check\_command* option in the service definitions that can be used to actively check the status of the service from the central monitoring server. Under normal circumstances, this check command is never executed (because active checks were disabled on a program-wide basis or for the specific services). When freshness checking is enabled, Nagios will run this command to actively check the status of the service *even if active checks are disabled on a program-wide or service-specific basis*.

If you are unable to define commands to actively check a service from the central monitoring host (or if turns out to be a major pain), you could simply define all your services with the *check\_command* option set to run a dummy script that returns a critical status. Here's an example... Let's assume you define a command called 'service-is-stale' and use that command name in the *check\_command* option of your services. Here's what the definition would look like...

```
define command{
    command_name    service-is-stale
    command_line    /usr/local/nagios/libexec/staleservice.sh
}
```

The **staleservice.sh** script in your */usr/local/nagios/libexec* directory might look something like this:

```
#!/bin/sh

/bin/echo "CRITICAL: Service results are stale!"

exit 2
```

When Nagios detects that the service results are stale and runs the **service-is-stale** command, the */usr/local/nagios/libexec/staleservice.sh* script is executed and the service will go into a critical state. This would likely cause notifications to be sent out, so you'll know that there's a problem.

## **Performing Host Checks**

At this point you know how to obtain service check results passively from distributed servers. This means that the central server is not actively checking services on its own. But what about host checks? You still need to do them, so how?

Since host checks usually compromise a small part of monitoring activity (they aren't done unless absolutely necessary), I'd recommend that you perform host checks actively from the central server. That means that you define host checks on the central server the same way that you do on the distributed servers (and the same way you would in a normal, non-distributed setup).

Passive host checks are available (read [here](#)), so you could use them in your distributed monitoring setup, but they suffer from a few problems. The biggest problem is that Nagios does not translate passive host check problem states (DOWN and UNREACHABLE) when they are processed. This means that if your monitoring servers have a different parent/child host structure (and they will, if you monitoring servers are in different locations), the central monitoring server will have an inaccurate view of host states.

If you do want to send passive host checks to a central server in your distributed monitoring setup, make sure:

- The central server has [passive host checks](#) enabled (required)
- The distributed server is configured to [obsess over hosts](#).
- The distributed server has an [ochp command](#) defined.

The ochp command, which is used for processing host check results, works in a similar manner to the oosp command, which is used for processing service check results (see documentation above). In order to make sure passive host check results are up to date, you'll want to enable [freshness checking](#) for hosts (similar to what is described above for services).

---



# Redundant and Failover Network Monitoring

---

## Introduction

This section describes a few scenarios for implementing redundant monitoring hosts on various types of network layouts. With redundant hosts, you can maintain the ability to monitor your network when the primary host that runs Nagios fails or when portions of your network become unreachable.

**Note:** If you are just learning how to use Nagios, I would suggest not trying to implement redundancy until you have becoming familiar with the [prerequisites](#) I've laid out. Redundancy is a relatively complicated issue to understand, and even more difficult to implement properly.

## Index

[Prerequisites](#)

[Sample scripts](#)

[Scenario 1 - Redundant monitoring](#)

[Scenario 2 - Failover monitoring](#)

### Prerequisites

Before you can even think about implementing redundancy with Nagios, you need to be familiar with the following...

- Implementing [event handlers](#) for hosts and services
- Issuing [external commands](#) to Nagios via shell scripts
- Executing plugins on remote hosts using either the [nrpe add-on](#) or some other method
- Checking the status of the Nagios process with the [check\\_nagios](#) plugin

### Sample Scripts

All of the sample scripts that I use in this documentation can be found in the *eventhandlers/* subdirectory of the Nagios distribution. You'll probably need to modify them to work on your system...

### Scenario 1 - Redundant Monitoring

## Introduction

This is an easy (and naive) method of implementing redundant monitoring hosts on your network and it will only protect against a limited number of failures. More complex setups are necessary in order to provide smarter redundancy, better redundancy across different network segments, etc.

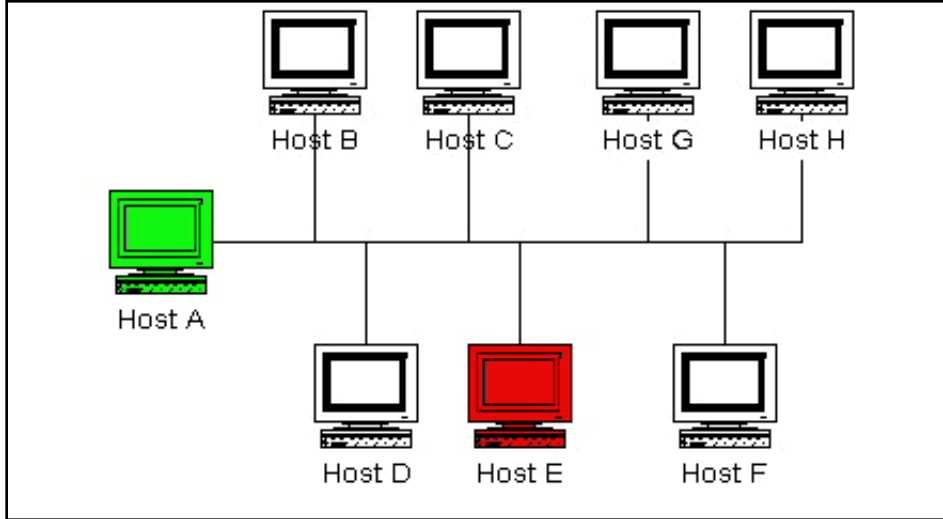
## Goals

The goal of this type of redundancy implementation is simple. Both the "master" and "slave" hosts monitor the same hosts and service on the network. Under normal circumstances only the "master" host will be sending out notifications to contacts about problems. We want the "slave" host running Nagios to take over the job of notifying contacts about problems if:

1. The "master" host that runs Nagios is down or..
2. The Nagios process on the "master" host stops running for some reason

## Network Layout Diagram

The diagram below shows a very simple network setup. For this scenario I will be assuming that hosts A and E are both running Nagios and are monitoring all the hosts shown. Host A will be considered the "master" host and host E will be considered the "slave" host.



## Initial Program Settings

The slave host (host E) has its initial [enable\\_notifications](#) directive disabled, thereby preventing it from sending out any host or service notifications. You also want to make sure that the slave host has its [check\\_external\\_commands](#) directive enabled. That was easy enough...

## Initial Configuration

Next we need to consider the differences between the [object configuration file\(s\)](#) on the master and slave hosts...

I will assume that you have the master host (host A) setup to monitor services on all hosts shown in the diagram above. The slave host (host E) should be setup to monitor the same services and hosts, with the following additions in the configuration file...

- The host definition for host A (in the host E configuration file) should have a host [event handler](#) defined. Lets say the name of the host event handler is **handle-master-host-event**.
- The configuration file on host E should have a service defined to check the status of the Nagios process on host A. Lets assume that you define this service check to run the *check\_nagios* plugin on host A (using NRPE, SSH, etc.).
- The service definition for the Nagios process check on host A should have an [event handler](#) defined. Lets say the name of the service event handler is **handle-master-proc-event**.

It is important to note that host A (the master host) has no knowledge of host E (the slave host). In this scenario it simply doesn't need to. Of course you may be monitoring services on host E from host A, but that has nothing to do with the implementation of redundancy...

## Event Handler Command Definitions

We need to stop for a minute and describe what the command definitions for the event handlers on the slave host look like. Here is an example...

```

define command{
    command_name    handle-master-host-event
    command_line    /usr/local/nagios/libexec/eventhandlers/handle-master-host-event $HOSTSTATE$ $HOSTSTATETYPE$
}

define command{
    command_name    handle-master-proc-event
    command_line    /usr/local/nagios/libexec/eventhandlers/handle-master-proc-event $SERVICESTATE$ $SERVICESTATETYPE$
}

```

This assumes that you have placed the event handler scripts in the */usr/local/nagios/libexec/eventhandlers* directory. You may place them anywhere you wish, but you'll need to modify the examples I've given here.

## **Event Handler Scripts**

Okay, now lets take a look at what the event handler scripts look like...

Host Event Handler (**handle-master-host-event**):

```

#!/bin/sh

# Only take action on hard host states...
case "$2" in
    HARD)
        case "$1" in
            DOWN)
                # The master host has gone down!
                # We should now become the master host and take
                # over the responsibilities of monitoring the
                # network, so enable notifications...
                /usr/local/nagios/libexec/eventhandlers/enable_notifications
                ;;
            UP)
                # The master host has recovered!
                # We should go back to being the slave host and
                # let the master host do the monitoring, so
                # disable notifications...
                /usr/local/nagios/libexec/eventhandlers/disable_notifications
                ;;
        esac
    ;;
esac
exit 0

```

Service Event Handler (**handle-master-proc-event**):

```

#!/bin/sh

# Only take action on hard service states...
case "$2" in
    HARD)
        case "$1" in
            CRITICAL)
                # The master Nagios process is not running!
                # We should now become the master host and
                # take over the responsibility of monitoring
                # the network, so enable notifications...
                /usr/local/nagios/libexec/eventhandlers/enable_notifications
                ;;
            WARNING)
            UNKNOWN)
                # The master Nagios process may or may not
                # be running.. We won't do anything here, but
                # to be on the safe side you may decide you
                # want the slave host to become the master in

```

```

        # these situations...
        ;;
    OK)
        # The master Nagios process running again!
        # We should go back to being the slave host,
        # so disable notifications...
        /usr/local/nagios/libexec/eventhandlers/disable_notifications
        ;;
    esac
    ;;
esac
exit 0

```

## What This Does For Us

The slave host (host E) initially has notifications disabled, so it won't send out any host or service notifications while the Nagios process on the master host (host A) is still running.

The Nagios process on the slave host (host E) becomes the master host when...

- The master host (host A) goes down and the *handle-master-host-event* host event handler is executed.
- The Nagios process on the master host (host A) stops running and the *handle-master-proc-event* service event handler is executed.

When the Nagios process on the slave host (host E) has notifications enabled, it will be able to send out notifications about any service or host problems or recoveries. At this point host E has effectively taken over the responsibility of notifying contacts of host and service problems!

The Nagios process on host E returns to being the slave host when...

- Host A recovers and the *handle-master-host-event* host event handler is executed.
- The Nagios process on host A recovers and the *handle-master-proc-event* service event handler is executed.

When the Nagios process on host E has notifications disabled, it will not send out notifications about any service or host problems or recoveries. At this point host E has handed over the responsibilities of notifying contacts of problems to the Nagios process on host A. Everything is now as it was when we first started!

## Time Lags

Redundancy in Nagios is by no means perfect. One of the more obvious problems is the lag time between the master host failing and the slave host taking over. This is affected by the following...

- The time between a failure of the master host and the first time the slave host detects a problem
- The time needed to verify that the master host really does have a problem (using service or host check retries on the slave host)
- The time between the execution of the event handler and the next time that Nagios checks for external commands

You can minimize this lag by...

- Ensuring that the Nagios process on host E (re)checks one or more services at a high frequency. This is done by using the *check\_interval* and *retry\_interval* arguments in each service definition.
- Ensuring that the number of host rechecks for host A (on host E) allow for fast detection of host problems. This is done by using the *max\_check\_attempts* argument in the host definition.
- Increase the frequency of *external command* checks on host E. This is done by modifying the

`command_check_interval` option in the main configuration file.

When Nagios recovers on the host A, there is also some lag time before host E returns to being a slave host. This is affected by the following...

- The time between a recovery of host A and the time the Nagios process on host E detects the recovery
- The time between the execution of the event handler on host B and the next time the Nagios process on host E checks for external commands

The exact lag times between the transfer of monitoring responsibilities will vary depending on how many services you have defined, the interval at which services are checked, and a lot of pure chance. At any rate, it's definitely better than nothing.

### **Special Cases**

Here is one thing you should be aware of... If host A goes down, host E will have notifications enabled and take over the responsibilities of notifying contacts of problems. When host A recovers, host E will have notifications disabled. If - when host A recovers - the Nagios process on host A does not start up properly, there will be a period of time when neither host is notifying contacts of problems! Fortunately, the service check logic in Nagios accounts for this. The next time the Nagios process on host E checks the status of the Nagios process on host A, it will find that it is not running. Host E will then have notifications enabled again and take over all responsibilities of notifying contacts of problems.

The exact amount of time that neither host is monitoring the network is hard to determine. Obviously, this period can be minimized by increasing the frequency of service checks (on host E) of the Nagios process on host A. The rest is up to pure chance, but the total "blackout" time shouldn't be too bad.

## **Scenario 2 - Failover Monitoring**

### **Introduction**

Failover monitoring is similar to, but slightly different than redundant monitoring (as discussed above in [scenario 1](#)).

### **Goals**

The basic goal of failover monitoring is to have the Nagios process on the slave host sit idle while the Nagios process on the master host is running. If the process on the master host stops running (or if the host goes down), the Nagios process on the slave host starts monitoring everything.

While the method described in [scenario 1](#) will allow you to continue receive notifications if the master monitoring hosts goes down, it does have some pitfalls. The biggest problem is that the slave host is monitoring the same hosts and servers as the master *at the same time as the master*! This can cause problems with excessive traffic and load on the machines being monitored if you have a lot of services defined. Here's how you can get around that problem...

### **Initial Program Settings**

Disable active service checks and notifications on the slave host using the [execute\\_service\\_checks](#) and [enable\\_notifications](#) directives. This will prevent the slave host from monitoring hosts and services and sending out notifications while the Nagios process on the master host is still up and running. Make sure you also have the [check\\_external\\_commands](#) directive enabled on the slave host.

## **Master Process Check**

Set up a cron job on the slave host that periodically (say every minute) runs a script that checks the status of the Nagios process on the master host (using the *check\_nrpe* plugin on the slave host and the [nrpe daemon](#) and *check\_nagios* plugin on the master host). The script should check the return code of the *check\_nrpe* plugin. If it returns a non-OK state, the script should send the appropriate commands to the [external command file](#) to enable both notifications and active service checks. If the plugin returns an OK state, the script should send commands to the external command file to disable both notifications and active checks.

By doing this you end up with only one process monitoring hosts and services at a time, which is much more efficient than monitoring everything twice.

Also of note, you *don't* need to define host and service handlers as mentioned in [scenario 1](#) because things are handled differently.

## **Additional Issues**

At this point, you have implemented a very basic failover monitoring setup. However, there is one more thing you should consider doing to make things work smoother.

The big problem with the way things have been setup thus far is the fact that the slave host doesn't have the current status of any services or hosts at the time it takes over the job of monitoring. One way to solve this problem is to enable the [ocsp command](#) on the master host and have it send all service check results to the slave host using the [nsca addon](#). The slave host will then have up-to-date status information for all services at the time it takes over the job of monitoring things. Since active service checks are not enabled on the slave host, it will not actively run any service checks. However, it will execute host checks if necessary. This means that both the master and slave hosts will be executing host checks as needed, which is not really a big deal since the majority of monitoring deals with service checks.

That's pretty much it as far as setup goes.

---

# Detection and Handling of State Flapping

## Introduction

Nagios supports optional detection of hosts and services that are "flapping". Flapping occurs when a service or host changes state too frequently, resulting in a storm of problem and recovery notifications. Flapping can be indicative of configuration problems (i.e. thresholds set too low) or real network problems.

Before I go any further, let me say that flapping detection has been a little difficult to implement. How exactly does one determine what "too frequently" means in regards to state changes for a particular host or service? When I first started looking into flap detection I tried to find some information on how flapping could/should be detected. After I couldn't find any, I decided to settle with what seemed to be a reasonable solution. The methods by which Nagios detects service and host state flapping are described below...

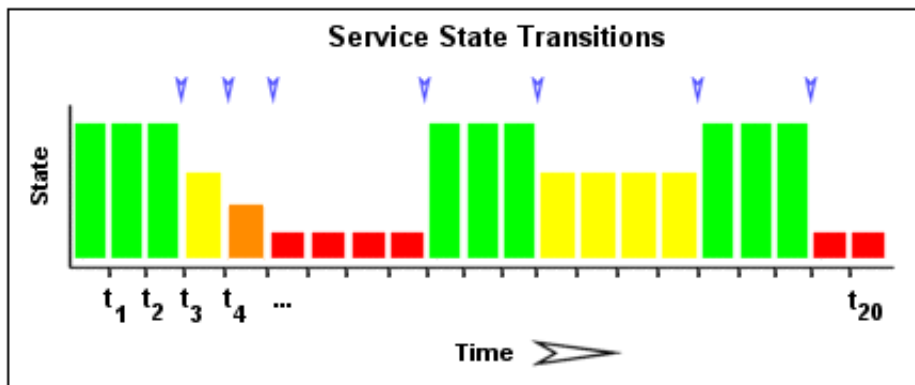
## Service Flap Detection

Whenever a service check is performed that results in a [hard state](#) or a soft recovery state, Nagios checks to see if the service has started or stopped flapping. It does this by storing the results of the last 21 checks of the service in an array. Older check results in the array are overwritten by newer check results.

The contents of the historical state array are examined (in order from oldest result to newest result) to determine the total percentage of change in state that has occurred during the last 21 service checks. A state change occurs when an archived state is different from the archived state that immediately precedes it in the array. Since we keep the results of the last 21 service checks in the array, there is a possibility of having 20 state changes.

Image 1 below shows a chronological array of service states. OK states are shown in green, WARNING states in yellow, CRITICAL states in red, and UNKNOWN states in orange. Blue arrows have been placed over periods of time where state changes occur.

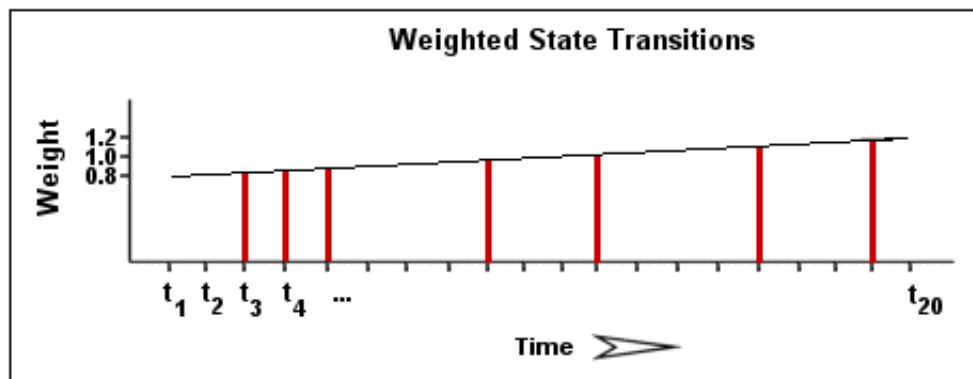
Image 1.



Services that rarely change between states will have a lower total percentage of change than those that do change between states a lot. Since flapping is associated with frequent state changes, we can use the calculated amount of change in state over a period of time (in this case, the last 21 service checks) to determine whether or not a service is flapping. That's not quite good enough though...

It stands to reason that newer state changes should carry more weight than older state changes, so we really need to recalculate the total percent change in state for the service on some sort of curve... To make things simple, I've decided to make the relationship between time and weight linear for calculation of percent state change. The flap detection routines are currently designed to make the newest possible state change carry 50% more weight than the oldest possible state change. Image 2 shows how more recent state changes are given more weight than older state changes when calculating the overall or total percent state change for a particular service. If you really want to see exactly how the weighted calculation is done, look at the code in `base/flapping.c`...

Image 2.



Let's look at a quick example of how flap detection is done. Image 1 above depicts the array of historical service check results for a particular service. The oldest result is on the left and the newest result is on the right. We see that in the example below there were a total of 7 state changes (at  $t_3, t_4, t_5, t_9, t_{12}, t_{16}$ , and  $t_{19}$ ). Without any weighting of the state changes over time, this would give us a total state change of 35% (7 state changes out of a possible 20 state changes). When the individual state changes are weighted relative to the time at which they occurred, the resulting total percent state change for the service is less than 35%. This makes sense since most of the state changes occurred earlier rather than later. Let's just say that the weighted percent of state change turned out to be 31%...

So what significance does the 31% state change have? Well, if the service was previously *not* flapping and 31% is *equal to or greater than* the value specified by the [high\\_service\\_flap\\_threshold](#) option in the service definition, Nagios considers the service to have just started flapping. If the service *was* previously flapping and 31% is *less than or equal to* the value specified by the [low\\_service\\_flap\\_threshold](#) value in the service definition, Nagios considers the service to have just stopped flapping. If either of those two conditions are not met, Nagios does nothing else with the service, since it is either not currently flapping or it is still flapping...

### Host Flap Detection

Host flap detection works in a similar manner to service flap detection, with one important difference: Nagios will attempt to check to see if a host is flapping whenever the status of the host is checked *and* whenever a service associated with that host is checked. Why is this done? Well, with services we know that the minimum amount of time between consecutive flap detection routines is going to be equal to the service check interval. With hosts, we don't have a check interval, since hosts are not monitored on a regular basis - they are only checked as necessary. A host will be checked for flapping if its state has changed since the last time the flap detection was performed for that host *or* if its state has not changed but at least  $x$  amount of time has passed since the flap detection was performed. The  $x$  amount of time is equal to the average check interval of all services associated with the host. That's the best method I could come up with for determining how often flap detection could be performed on a host...



Just as with services, Nagios stores the results of the last 21 of these host checks in an array for the flap detection logic. State changes are weighted based on the time at which they occurred, and the total percent change in state is calculated in the same manner that it is in the service flapping logic.

If a host was previously *not* flapping and its total computed state change percentage is *equal to or greater than* the value specified by the [high\\_host\\_flap\\_threshold](#) option, Nagios considers the host to have just started flapping. If the host *was* previously flapping and its total computed state change percentage is *less than or equal to* the value specified by the [low\\_host\\_flap\\_threshold](#) value, Nagios considers the host to have just stopped flapping. If either of those two conditions are not met, Nagios does nothing else with the host, since it is either not currently flapping or it is still flapping...

### **Host- and Service-Specific Flap Detection Thresholds**

If you're using the [template-based object definition files](#), you can specify host- and service-specific flap detection thresholds by adding **low\_flap\_threshold** and **high\_flap\_threshold** directives to individual host and service definitions. If these directives are *not* present in the host or service definitions, the global host and service flap detection thresholds will be used.

On a similar note, you can also enable/disable flap detection for specific hosts and services by using the **flap\_detection\_enabled** directive in each object definition. Note that flap detection must be enabled on a program-wide basis (using the [enable\\_flap\\_detection](#) directive in the main config file) in order for any host or service to have flap detection enabled.

### **Flap Handling**

When a service or host is first detected as flapping, Nagios will do three things:

1. Log a message indicating that the service or host is flapping
2. Add a non-persistent comment to the host or service indicating that it is flapping
3. Suppress notifications for the service or host (this is one of the filters in the [notification logic](#))

When a service or host stops flapping, Nagios will do the following:

1. Log a message indicating that the service or host has stopped flapping
  2. Delete the comment that was originally added to the service or host when it started flapping
  3. Remove the block on notifications for the service or host (notifications will still be bound to the normal [notification logic](#))
-

# Service Check Parallelization

---

## Introduction

One of the features of Nagios is its ability to execute service checks in parallel. This documentation will attempt to explain in detail what that means and how it affects services that you have defined.

## How The Parallelization Works

Before I can explain how the service check parallelization works, you first have to understand a bit about how Nagios schedules events. All internal events in Nagios (i.e. log file rotations, external command checks, service checks, etc.) are placed in an event queue. Each item in the event queue has a time at which it is scheduled to be executed. Nagios does its best to ensure that all events get executed when they should, although events may fall behind schedule if Nagios is busy doing other things.

Service checks are one type of event that get scheduled in Nagios' event queue. When it comes time for a service check to be executed, Nagios will kick off another process (using a call to `fork()`) to go out and run the service check (i.e. a plugin of some sort). Nagios does *not*, however, wait for the service check to finish! Instead, Nagios will immediately go back to servicing other events that reside in the event queue...

So what happens when the service check finishes executing? Well, the process that was started by Nagios to run the service check sends a message back to Nagios containing the results of the service check. It is then up to Nagios to check for and process the results of that service check when it gets a chance.

In order for Nagios to actually do any monitoring, it must process the results of service checks that have finished executing. This is done via a service check "reaper" process. Service "reapers" are another type of event that get scheduled in Nagios' event queue. The frequency of these "reaper" events is determined by the [service\\_reaper\\_frequency](#) option in the main configuration file. When a "reaper" event is executed, it will check for any messages that contain the result of service checks that have finished executing. These service check results are then handled by the core service monitoring logic. From there Nagios determines whether or not hosts should be checked, notifications should be sent out, etc. When the service check results have been processed, Nagios will reschedule the next check of the service and place it in the event queue for later execution. That completes the service check/monitoring cycle!

For those of you who really want to know, but haven't looked at the code, Nagios uses message queues to handle communication between Nagios and the process that actually runs the service check...

## Potential Gotchas...

You should realize that there are potential drawbacks to having service checks parallelized. Since more than one service check may be running at the same time, they have may interfere with one another. You'll have to evaluate what types of service checks you're running and take appropriate steps to guard against any unfriendly outcomes. This is particularly important if you have more than one service check that accesses any hardware (like a modem). Also, if two or more service checks connect to daemon on a remote host to check some information, make sure that daemon can handle multiple simultaneous connections.

Fortunately, there are some things you can do to protect against problems with having some types of service checks "collide"...

1. The easiest thing you can do to prevent service check collisions to to use the [service\\_interleave\\_factor](#) variable. Interleaving services will help to reduce the load imposed upon remote hosts by service checks. Set the variable to use "smart" interleave factor calculation and then adjust it manually if you find it necessary to do so.
2. The second thing you can do is to set the *max\_check\_attempts* argument in each service definition to something greater than one. If the service check does happen to collide with another running check, Nagios will retry the service check *max\_check\_attempts-1* times before notifying anyone of a problem.
3. You could try is to implement some kind of "back-off and retry" logic in the actual service check code, although you may find it difficult or too time-consuming
4. If all else fails you can effectively prevent service checks from being parallelized by setting the [max\\_concurrent\\_checks](#) option to 1. This will allow only one service to be checked at a time, so it isn't a spectacular solution. If there is enough demand, I will add an option to the service definitions which will allow you to specify on a per-service basis whether or not a service check can be parallelized. If there isn't enough demand, I won't...

One other thing to note is the effect that parallelization of service checks can have on system resources on the machine that runs Nagios. Running a lot of service checks in parallel can be taxing on the CPU and memory. The [inter\\_check\\_delay\\_method](#) will attempt to minimize the load imposed on your machine by spreading the checks out evenly over time (if you use the "smart" method), but it isn't a surefire solution. In order to have some control over how many service checks can be run at any given time, use the [max\\_concurrent\\_checks](#) variable. You'll have to tweak this value based on the total number of services you check, the system resources you have available (CPU speed, memory, etc.), and other processes which are running on your machine. For more information on how to tweak the *max\_concurrent\_checks* variable for your setup, read the documentation on [check scheduling](#).

### **What Isn't Parallelized**

It is important to remember that only the *execution* of service checks has been parallelized. There is good reason for this - other things cannot be parallelized in a very safe or sane manner. In particular, event handlers, contact notifications, processing of service checks, and host checks are *not* parallelized. Here's why...

*Event handlers* are not parallelized because of what they are designed to do. Much of the power of event handlers comes from the ability to do proactive problem resolution. An example of this is restarting the web server when the HTTP service on the local machine is detected as being down. In order to prevent more than one event handler from trying to "fix" problems in parallel (without any knowledge of what each other is doing), I have decided to not parallelize them.

*Contact notifications* are not parallelized because of potential notification methods you may be using. If, for example, a contact notification uses a modem to dial out and send a message to your pager, it requires exclusive access to the modem while the notification is in progress. If two or more such notifications were being executed in parallel, all but one would fail because the others could not get access to the modem. There are ways to get around this, like providing some kind of "back-off and retry" method in the notification script, but I've decided not to rely on users having implemented this type of feature in their scripts. One quick note - if you have service checks which use a modem, make sure that any notification scripts that dial out have some method of retrying access to the modem. This is necessary because a service check may be running at the same time a notification is!

*Processing of service check results* has not been parallelized. This has been done to prevent situations where multiple notifications about host problems or recoveries may be sent out if a host goes down, becomes unreachable, or recovers.

---

# Notification Escalations

---

## Introduction

Nagios supports *optional* escalation of contact notifications for hosts and services. I'll explain quickly how they work, although they should be fairly self-explanatory...

## Service Notification Escalations

Escalation of service notifications is accomplished by defining [service escalations](#) in your [object configuration file](#). Service escalation definitions are used to escalate notifications for a particular service.

## Host Notification Escalations

Escalation of host notifications is accomplished by defining [host escalations](#) in your [object configuration file](#). The examples I provide below all use service escalation definitions, but host escalations work the same way (except for the fact that they are used for host notifications and not service notifications).

## When Are Notifications Escalated?

Notifications are escalated *if and only if* one or more escalation definitions matches the current notification that is being sent out. If a host or service notification *does not* have any valid escalation definitions that applies to it, the contact group(s) specified in either the host group or service definition will be used for the notification. Look at the example below:

```
define serviceescalation{
    host_name           webserver
    service_description HTTP
    first_notification   3
    last_notification    5
    notification_interval 90
    contact_groups       nt-admins,managers
}

define serviceescalation{
    host_name           webserver
    service_description HTTP
    first_notification   6
    last_notification    10
    notification_interval 60
    contact_groups       nt-admins,managers,everyone
}
```

Notice that there are "holes" in the notification escalation definitions. In particular, notifications 1 and 2 are not handled by the escalations, nor are any notifications beyond 10. For the first and second notification, as well as all notifications beyond the tenth one, the *default* contact groups specified in the service definition are used. For all the examples I'll be using, I'll be assuming that the default contact groups for the service definition is called *nt-admins*.

## Contact Groups

When defining notification escalations, it is important to keep in mind that any contact groups that were members of "lower" escalations (i.e. those with lower notification number ranges) should also be included in "higher" escalation definitions. This should be done to ensure that anyone who gets notified of a problem *continues* to get notified as the problem is escalated. Example:

```

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     3
    last_notification      5
    notification_interval  90
    contact_groups         nt-admins,managers
}

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     6
    last_notification      0
    notification_interval  60
    contact_groups         nt-admins,managers,everyone
}

```

The first (or "lowest") escalation level includes both the *nt-admins* and *managers* contact groups. The last (or "highest") escalation level includes the *nt-admins*, *managers*, and *everyone* contact groups. Notice that the *nt-admins* contact group is included in both escalation definitions. This is done so that they continue to get paged if there are still problems after the first two service notifications are sent out. The *managers* contact group first appears in the "lower" escalation definition - they are first notified when the third problem notification gets sent out. We want the *managers* group to continue to be notified if the problem continues past five notifications, so they are also included in the "higher" escalation definition.

### **Overlapping Escalation Ranges**

Notification escalation definitions can have notification ranges that overlap. Take the following example:

```

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     3
    last_notification      5
    notification_interval  20
    contact_groups         nt-admins,managers
}

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     4
    last_notification      0
    notification_interval  30
    contact_groups         on-call-support
}

```

In the example above:

- The *nt-admins* and *managers* contact groups get notified on the third notification
- All three contact groups get notified on the fourth and fifth notifications
- Only the *on-call-support* contact group gets notified on the sixth (or higher) notification

### **Recovery Notifications**

Recovery notifications are slightly different than problem notifications when it comes to escalations. Take the following example:

```

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     3
    last_notification      5
    notification_interval  20
    contact_groups         nt-admins,managers
}

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     4
    last_notification      0
    notification_interval  30
    contact_groups         on-call-support
}

```

If, after three problem notifications, a recovery notification is sent out for the service, who gets notified? The recovery is actually the fourth notification that gets sent out. However, the escalation code is smart enough to realize that only those people who were notified about the problem on the third notification should be notified about the recovery. In this case, the *nt-admins* and *managers* contact groups would be notified of the recovery.

### **Notification Intervals**

You can change the frequency at which escalated notifications are sent out for a particular host or service by using the *notification\_interval* option of the hostgroup or service escalation definition. Example:

```

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     3
    last_notification      5
    notification_interval  45
    contact_groups         nt-admins,managers
}

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     6
    last_notification      0
    notification_interval  60
    contact_groups         nt-admins,managers,everyone
}

```

In this example we see that the default notification interval for the services is 240 minutes (this is the value in the service definition). When the service notification is escalated on the 3rd, 4th, and 5th notifications, an interval of 45 minutes will be used between notifications. On the 6th and subsequent notifications, the notification interval will be 60 minutes, as specified in the second escalation definition.

Since it is possible to have overlapping escalation definitions for a particular hostgroup or service, and the fact that a host can be a member of multiple hostgroups, Nagios has to make a decision on what to do as far as the notification interval is concerned when escalation definitions overlap. In any case where there are multiple valid escalation definitions for a particular notification, Nagios will choose the smallest notification interval. Take the following example:

```

define serviceescalation{
    host_name             webserver
    service_description    HTTP
    first_notification     3

```

```

        last_notification      5
        notification_interval  45
        contact_groups        nt-admins,managers
    }

define serviceescalation{
    host_name                webserver
    service_description      HTTP
    first_notification       4
    last_notification        0
    notification_interval    60
    contact_groups          nt-admins,managers,everyone
}

```

We see that the two escalation definitions overlap on the 4th and 5th notifications. For these notifications, Nagios will use a notification interval of 45 minutes, since it is the smallest interval present in any valid escalation definitions for those notifications.

One last note about notification intervals deals with intervals of 0. An interval of 0 means that Nagios should only sent a notification out for the first valid notification during that escalation definition. All subsequent notifications for the hostgroup or service will be suppressed. Take this example:

```

define serviceescalation{
    host_name                webserver
    service_description      HTTP
    first_notification       3
    last_notification        5
    notification_interval    45
    contact_groups          nt-admins,managers
}

define serviceescalation{
    host_name                webserver
    service_description      HTTP
    first_notification       4
    last_notification        6
    notification_interval    0
    contact_groups          nt-admins,managers,everyone
}

define serviceescalation{
    host_name                webserver
    service_description      HTTP
    first_notification       7
    last_notification        0
    notification_interval    30
    contact_groups          nt-admins,managers
}

```

In the example above, the maximum number of problem notifications that could be sent out about the service would be four. This is because the notification interval of 0 in the second escalation definition indicates that only one notification should be sent out (starting with and including the 4th notification) and all subsequent notifications should be repressed. Because of this, the third service escalation definition has no effect whatsoever, as there will never be more than four notifications.

### **Time Period Restrictions**

Under normal circumstances, escalations can be used at any time that a notification could normally be sent out for the service. This "notification time window" is determined by the *notification\_period* directive in the [service definition](#).

You can optionally restrict escalations so that they are only used during specific time periods by using the *escalation\_period* directive in the [service escalation](#) definition. If you use the *escalation\_period* directive to specify a [timeperiod](#) during which the escalation can be used, the escalation will only be used during that time. If you do not specify any *escalation\_period* directive, the escalation can be used at any time within the "notification time window" for the service.

Note that the notification is still subject to the normal time restrictions imposed by the *notification\_period* directive in the service escalation, so the timeperiod you specify in the escalation should be a subset of that larger "notification time window".

### **State Restrictions**

If you would like to restrict the escalation definition so that it is only used when the service is in a particular state, you can use the *escalation\_options* directive in the [service escalation](#) definition. If you do not use the *escalation\_options* directive, the escalation can be used when the service is in any state.

---



# Monitoring Service and Host Clusters

---

## Introduction

Several people have asked how to go about monitoring clusters of hosts or services, so I decided to write up a little documentation on how to do this. Its fairly straightforward, so hopefully you find things easy to understand...

First off, we need to define what we mean by a "cluster". The simplest way to understand this is with an example. Let's say that your organization has five hosts which provide redundant DNS services to your organization. If one of them fails, its not a major catastrophe because the remaining servers will continue to provide name resolution services. If you're concerned with monitoring the availability of DNS service to your organization, you will want to monitor five DNS servers. This is what I consider to be a *service* cluster. The service cluster consists of five separate DNS services that you are monitoring. Although you do want to monitor each individual service, your main concern is with the overall status of the DNS service cluster, rather than the availability of any one particular service.

If your organization has a group of hosts that provide a high-availability (clustering) solution, I would consider those to be a *host* cluster. If one particular host fails, another will step in to take over all the duties of the failed server. As a side note, check out the [High-Availability Linux Project](#) for information on providing host and service redundancy with Linux.

## Plan of Attack

There are several ways you could potentially monitor service or host clusters. I'll describe the method that I believe to be the easiest. Monitoring service or host clusters involves two things:

- Monitoring individual cluster elements
- Monitoring the cluster as a collective entity

Monitoring individual host or service cluster elements is easier than you think. In fact, you're probably already doing it. For service clusters, just make sure that you are monitoring each service element of the cluster. If you've got a cluster of five DNS servers, make sure you have five separate service definitions (probably using the `check_dns` plugin). For host clusters, make sure you have configured appropriate host definitions for each member of the cluster (you'll also have to define at least one service to be monitored for each of the hosts). **Important:** You're going to want to disable notifications for the individual cluster elements (host or service definitions). Even though no notifications will be sent about the individual elements, you'll still get a visual display of the individual host or service status in the [status CGI](#). This will be useful for pinpointing the source of problems within the cluster in the future.

Monitoring the overall cluster can be done by using the previously cached results of cluster elements. Although you could re-check all elements of the cluster to determine the cluster's status, why waste bandwidth and resources when you already have the results cached? Where are the results cached? Cached results for cluster elements can be found in the [status file](#) (assuming you are monitoring each element). The `check_cluster` plugin is designed specifically for checking cached host and service states in the status file. **Important:** Although you didn't enable notifications for individual elements of the cluster, you will want them enabled for the overall cluster status check.

## Using the `check_cluster` Plugin

The `check_cluster` plugin is designed to report the overall status of a host or service cluster by checking the status information of each individual host or service cluster elements.

More to come... The *check\_cluster* plugin can be found in the contrib directory of the Nagios Plugins release at <http://sourceforge.net/projects/nagiosplug/>.

## Monitoring Service Clusters

Let's say you have three DNS servers that provide redundant services on your network. First off, you need to be monitoring each of these DNS servers separately before you can monitor them as a cluster. I'll assume that you already have three separate services (all called "DNS Service") associated with your DNS hosts (called "host1", "host2" and "host3").

In order to monitor the services as a cluster, you'll need to create a new "cluster" service. However, before you do that, make sure you have a service cluster check command configured. Let's assume that you have a command called *check\_service\_cluster* defined as follows:

```
define command{
    command_name    check_service_cluster
    command_line    /usr/local/nagios/libexec/check_cluster --service -l $ARG1$ -w $ARG2$ -c $ARG3$ -d $ARG4$
}
```

Now you'll need to create the "cluster" service and use the *check\_service\_cluster* command you just created as the cluster's check command. The example below gives an example of how to do this. The example below will generate a CRITICAL alert if 2 or more services in the cluster are in a non-OK state, and a WARNING alert if only 1 of the services is in a non-OK state. If all the individual service members of the cluster are OK, the cluster check will return an OK state as well.

```
define service{
    ...
    check_command    check_service_cluster!"DNS Cluster"!1!2!$SERVICEID:host1:DNS Service$, $SERVICEID
    ...
}
```

It is important to notice that we are passing a comma-delimited list of *on-demand* service state [macros](#) to the \$ARG4\$ macro in the cluster check command. That's important! Nagios will fill those on-demand macros in with the current service state IDs (numerical values, rather than text strings) of the individual members of the cluster.

## Monitoring Host Clusters

Monitoring host clusters is very similar to monitoring service clusters. Obviously, the main difference is that the cluster members are hosts and not services. In order to monitor the status of a host cluster, you must define a service that uses the *check\_cluster* plugin. The service should *not* be associated with any of the hosts in the cluster, as this will cause problems with notifications for the cluster if that host goes down. A good idea might be to associate the service with the host that Nagios is running on. After all, if the host that Nagios is running on goes down, then Nagios isn't running anymore, so there isn't anything you can do as far as monitoring (unless you've setup [redundant monitoring hosts](#))...

Anyway, let's assume that you have a *check\_host\_cluster* command defined as follows:

```
define command{
    command_name    check_host_cluster
    command_line    /usr/local/nagios/libexec/check_cluster --host -l $ARG1$ -w $ARG2$ -c $ARG3$ -d $ARG4$
}
```

Let's say you have three hosts (named "host1", "host2" and "host3") in the host cluster. If you want Nagios to generate a warning alert if one host in the cluster is not UP or a critical alert if two or more hosts are not UP, the the service you define to monitor the host cluster might look something like this:

```
define service{
    ...
    check_command    check_host_cluster!"Super Host Cluster"!1!2!$HOSTSTATEID:host1$, $HOSTSTATEID:host2$, ...
    ...
}
```

It is important to notice that we are passing a comma-delimited list of *on-demand* host state [macros](#) to the \$ARG4\$ macro in the cluster check command. That's important! Nagios will fill those on-demand macros in with the current host state IDs (numerical values, rather than text strings) of the individual members of the cluster.

That's it! Nagios will periodically check the status of the host cluster and send notifications to you when its status is degraded (assuming you've enabled notification for the service). Note that for the host definitions of each cluster member, you will most likely want to disable notifications when the host goes down . Remember that you don't care as much about the status of any individual host as you do the overall status of the cluster. Depending on your network layout and what you're trying to accomplish, you may wish to leave notifications for unreachable states enabled for the host definitions.

---

# Host and Service Dependencies

---

## Introduction

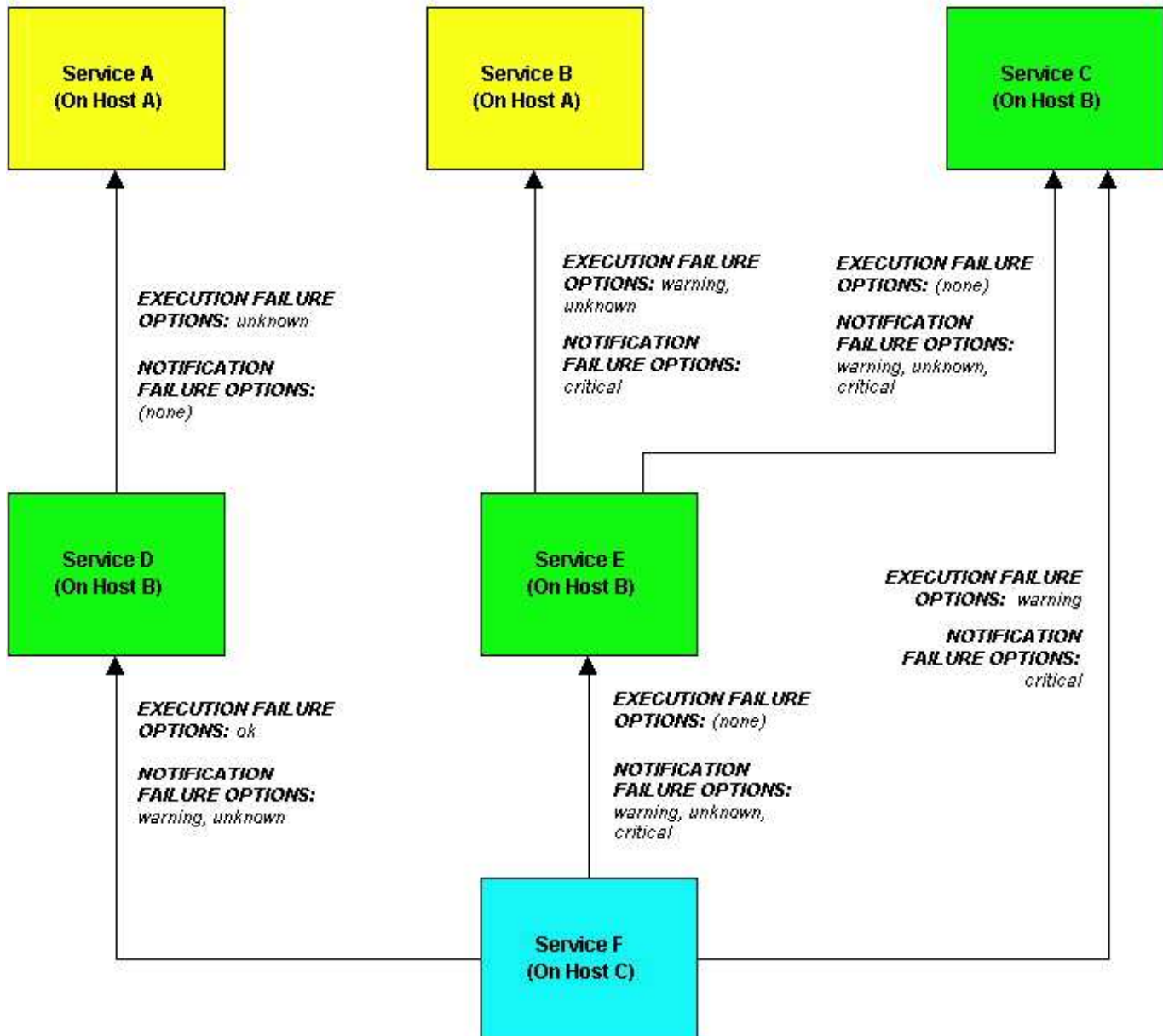
Service and host dependencies are an *advanced* feature that allow you to control the behavior of hosts and services based on the status of one or more other hosts or services. I'll explain how dependencies work, along with the differences between host and service dependencies.

## Service Dependencies Overview

The image below shows an example logical layout of service dependencies. There are a few things you should notice:

1. A service can be dependent on one or more other services
2. A service can be dependent on services which are not associated with the same host
3. Service dependencies are not inherited (unless specifically configured to)
4. Service dependencies can be used to cause service execution and service notifications to be suppressed under different circumstances (OK, WARNING, UNKNOWN, and/or CRITICAL states)

# Service Dependencies



## Defining Service Dependencies

First, the basics. You create service dependencies by adding service dependency definitions in your [object config file\(s\)](#). In each definition you specify the *dependent* service, the service you are *depending on*, and the criteria (if any) that cause the execution and notification dependencies to fail (these are described later).

You can create several dependencies for a given service, but you must add a separate service dependency definition for each dependency you create.

In the image above, the dependency definitions for *Service F* on *Host C* would be defined as follows:

```
define servicedependency{
  host_name           Host B
  service_description Service D
  dependent_host_name Host C
```

```

        dependent_service_description    Service F
        execution_failure_criteria       o
        notification_failure_criteria    w,u
    }

define servicedependency{
    host_name                Host B
    service_description       Service E
    dependent_host_name       Host C
    dependent_service_description    Service F
    execution_failure_criteria    n
    notification_failure_criteria  w,u,c
}

define servicedependency{
    host_name                Host B
    service_description       Service C
    dependent_host_name       Host C
    dependent_service_description    Service F
    execution_failure_criteria    w
    notification_failure_criteria  c
}

```

The other dependency definitions shown in the image above would be defined as follows:

```

define servicedependency{
    host_name                Host A
    service_description       Service A
    dependent_host_name       Host B
    dependent_service_description    Service D
    execution_failure_criteria    u
    notification_failure_criteria  n
}

define servicedependency{
    host_name                Host A
    service_description       Service B
    dependent_host_name       Host B
    dependent_service_description    Service E
    execution_failure_criteria    w,u
    notification_failure_criteria  c
}

define servicedependency{
    host_name                Host B
    service_description       Service C
    dependent_host_name       Host B
    dependent_service_description    Service E
    execution_failure_criteria    n
    notification_failure_criteria  w,u,c
}

```

## **How Service Dependencies Are Tested**

Before Nagios executes a service check or sends notifications out for a service, it will check to see if the service has any dependencies. If it doesn't have any dependencies, the check is executed or the notification is sent out as it normally would be. If the service *does* have one or more dependencies, Nagios will check each dependency entry as follows:

1. Nagios gets the current status\* of the service that is being *depended upon*.
2. Nagios compares the current status of the service that is being *depended upon* against either the execution or notification failure options in the dependency definition (whichever one is relevant at the time).

3. If the current status of the service that is being *depended upon* matches one of the failure options, the dependency is said to have failed and Nagios will break out of the dependency check loop.
4. If the current state of the service that is being *depended upon* does not match any of the failure options for the dependency entry, the dependency is said to have passed and Nagios will go on and check the next dependency entry.

This cycle continues until either all dependencies for the service have been checked or until one dependency check fails.

\*One important thing to note is that by default, Nagios will use the most current [hard](#) state of the service(s) that is/are being depended upon when it does the dependency checks. If you want Nagios to use the most current state of the services (regardless of whether its a soft or hard state), enable the [soft\\_service\\_dependencies](#) option.

### **Execution Dependencies**

Execution dependencies are used to restrict when *active* checks of a service can be performed. Passive checks are not restricted by execution dependencies.

If *all* of the execution dependency tests for the service *passed*, Nagios will execute the check of the service as it normally would. If even just one of the execution dependencies for a service fails, Nagios will temporarily prevent the execution of checks for that (dependent) service. At some point in the future the execution dependency tests for the service may all pass. If this happens, Nagios will start checking the service again as it normally would. More information on the check scheduling logic can be found [here](#).

In the example above, **Service E** would have failed execution dependencies if **Service B** is in a WARNING or UNKNOWN state. If this was the case, the service check would not be performed and the check would be scheduled for (potential) execution at a later time.

### **Notification Dependencies**

If *all* of the notification dependency tests for the service *passed*, Nagios will send notifications out for the service as it normally would. If even just one of the notification dependencies for a service fails, Nagios will temporarily repress notifications for that (dependent) service. At some point in the future the notification dependency tests for the service may all pass. If this happens, Nagios will start sending out notifications again as it normally would for the service. More information on the notification logic can be found [here](#).

In the example above, **Service F** would have failed notification dependencies if **Service C** is in a CRITICAL state, *and/or* **Service D** is in a WARNING or UNKNOWN state, *and/or* if **Service E** is in a WARNING, UNKNOWN, or CRITICAL state. If this were the case, notifications for the service would not be sent out.

### **Dependency Inheritance**

As mentioned before, service dependencies are *not* inherited by default. In the example above you can see that Service F is dependent on Service E. However, it does not automatically inherit Service E's dependencies on Service B and Service C. In order to make Service F dependent on Service C we had to add another service dependency definition. There is no dependency definition for Service B, so Service F is *not* dependent on Service B.

If you *do* wish to make service dependencies inheritable, you must use the *inherits\_parent* directive in the [service dependency](#) definition. When this directive is enabled, it indicates that the dependency inherits dependencies of the service *that is being depended upon* (also referred to as the master service). In other words, if the master service is dependent upon other services and any one of those dependencies fail, this dependency will also fail.

In the example above, imagine that you want to add a new dependency for service F to make it dependent on service A. You could create a new dependency definition that specified service F as the *dependent* service and service A as being the *master* service (i.e. the service *that is being dependend on*). You could alternatively modify the dependency definition for services D and F to look like this:

```
define servicedependency{
    host_name           Host B
    service_description  Service D
    dependent_host_name  Host C
    dependent_service_description  Service F
    execution_failure_criteria  o
    notification_failure_criteria  n
    inherits_parent      1
}
```

Since the *inherits\_parent* directive is enabled, the dependency between services A and D will be tested when the dependency between services F and D are being tested.

Dependencies can have multiple levels of inheritance. If the dependency definition between A and D had its *inherits\_parent* directive enable and service A was dependent on some other service (let's call it service G), the service F would be dependent on services D, A, and G (each with potentially different criteria).

### **Host Dependencies**

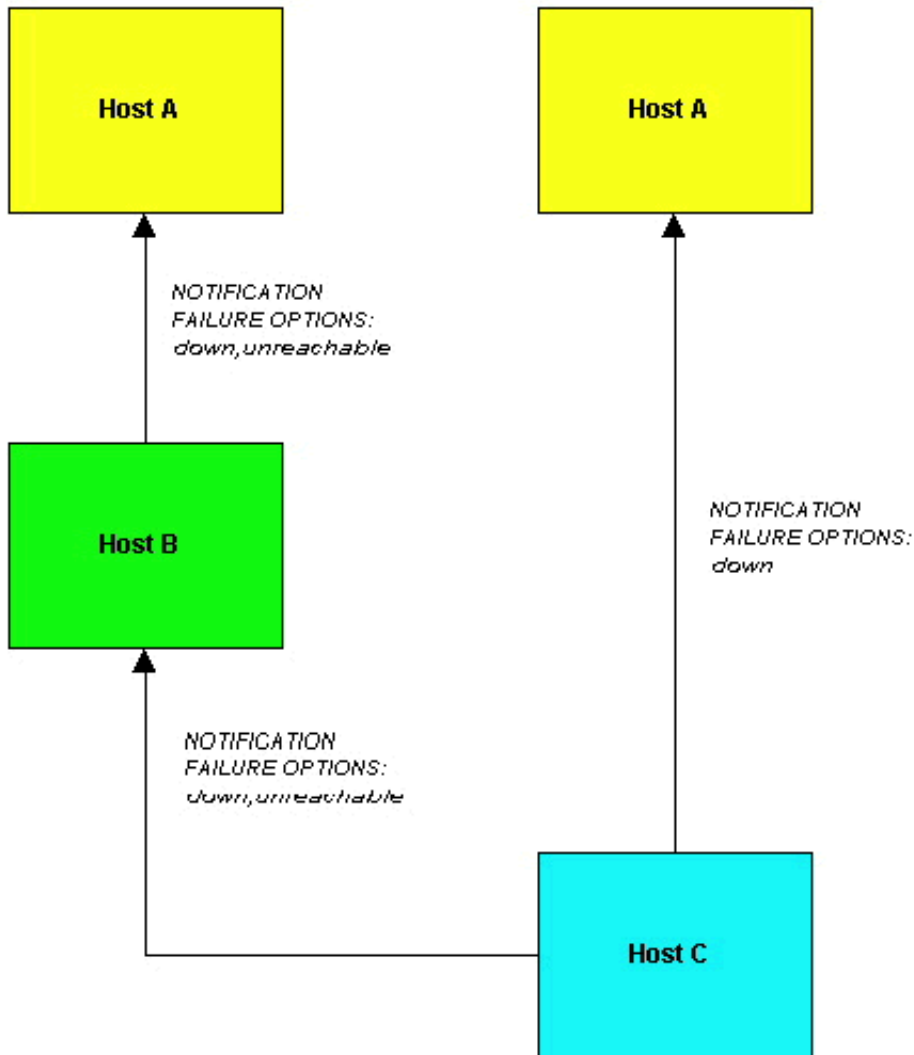
As you'd probably expect, host dependencies work in a similiar fashion to service dependencies. The big difference is that they're for hosts, not services. Another difference is that host dependencies only work for suppressing host notifications, not host checks.

**BEWARE!** Do not confuse host dependencies with parent/child host relationships. You should be using parent/child host relationships (defined with the *parents* directive in [host](#) definitions) for most cases, rather than host dependencies.

The image below shows an example of the logical layout of host dependencies.



# Host Dependencies



In the image above, the dependency definitions for *Host C* would be defined as follows:

```
define hostdependency{
    host_name           Host A
    dependent_host_name Host C
    notification_failure_criteria d
}

define hostdependency{
    host_name           Host B
    dependent_host_name Host C
    notification_failure_criteria d,u
}
```

As with service dependencies, host dependencies are not inherited. In the example image you can see that Host C does not inherit the host dependencies of Host B. In order for Host C to be dependent on Host A, a new host dependency definition must be defined.

Host notification dependencies work in a similar manner to service notification dependencies. If *all* of the notification dependency tests for the host *pass*, Nagios will send notifications out for the host as it normally would. If even just one of the notification dependencies for a host fails, Nagios will

temporarily repress notifications for that (dependent) host. At some point in the future the notification dependency tests for the host may all pass. If this happens, Nagios will start sending out notifications again as it normally would for the host. More information on the notification logic can be found [here](#).

**NOTE:** Host execution dependencies work in a similiar manner to service execution dependencies. However, they only have an affect on *regularly scheduled host checks*. On-demand host checks are not affected by host execution dependencies.

---

# State Stalking

---

## Introduction

State "stalking" is a feature which is probably not going to be used by most users. When enabled, it allows you to log changes in service and host checks even if the state of the host or service does not change. When stalking is enabled for a particular host or service, Nagios will watch that service very carefully and log any changes it sees. As you'll see, it can be very helpful to you in later analysis of the log files.

## How Does It Work?

Under normal circumstances, the result of a host or service check is only logged if the host or service has changed state since it was last checked. There are a few exceptions to this, but for the most part, that's the rule.

If you enable stalking for one or more states of a particular host or service, Nagios will log the results of the host or service check if the output from the check differs from the output from the previous check. Take the following example of eight consecutive checks of a service:

| Service Check #: | Service State: | Service Check Output:                                                          |
|------------------|----------------|--------------------------------------------------------------------------------|
| x                | OK             | RAID array optimal                                                             |
| x+1              | OK             | RAID array optimal                                                             |
| x+2              | WARNING        | RAID array degraded (1 drive bad, 1 hot spare rebuilding)                      |
| x+3              | CRITICAL       | RAID array degraded (2 drives bad, 1 hot spare online, 1 hot spare rebuilding) |
| x+4              | CRITICAL       | RAID array degraded (3 drives bad, 2 hot spares online)                        |
| x+5              | CRITICAL       | RAID array failed                                                              |
| x+6              | CRITICAL       | RAID array failed                                                              |
| x+7              | CRITICAL       | RAID array failed                                                              |

Given this sequence of checks, you would normally only see two log entries for this catastrophe. The first one would occur at service check x+2 when the service changed from an OK state to a WARNING state. The second log entry would occur at service check x+3 when the service changed from a WARNING state to a CRITICAL state.

For whatever reason, you may like to have the complete history of this catastrophe in your log files. Perhaps to help explain to your manager how quickly the situation got out of control, perhaps just to laugh at over a couple of drinks at the local pub, whatever...

Well, if you had enabled stalking of this service for CRITICAL states, you would have events at x+4 and x+5 logged in addition to the events at x+2 and x+3. Why is this? With state stalking enabled, Nagios would have examined the output from each service check to see if it differed from the output of the previous check. If the output differed and the state of the service didn't change between the two checks, the result of the newer service check would get logged.

A similar example of stalking might be on a service that checks your web server. If the `check_http` plugin first returns a WARNING state because of a 404 error and on subsequent checks returns a WARNING state because of a particular pattern not being found, you might want to know that. If you didn't enable state stalking for WARNING states of the service, only the first WARNING state event (the 404 error) would be logged and you wouldn't have any idea (looking back in the archived logs) that future problems were not due to a 404, but rather a missing pattern in the returned web page.

### **Should I Enable Stalking?**

First, you must decide if you have a real need to analyze archived log data to find the exact cause of a problem. You may decide you need this feature for some hosts or services, but not for all. You may also find that you only have a need to enable stalking for some host or service states, rather than all of them. For example, you may decide to enable stalking for WARNING and CRITICAL states of a service, but not for OK and UNKNOWN states.

The decision to to enable state stalking for a particular host or service will also depend on the plugin that you use to check that host or service. If the plugin always returns the same text output for a particular state, there is no reason to enable stalking for that state.

### **How Do I Enable Stalking?**

You can enable state stalking for hosts and services by using the *stalking\_options* directive in [host and service definitions](#).

### **Caveats**

You should be aware that there are some potential pitfalls with enabling stalking. These all relate to the reporting functions found in various [CGIs](#) (histogram, alert summary, etc.). Because state stalking will cause additional alert entries to be logged, the data produced by the reports will show evidence of inflated numbers of alerts.

As a general rule, I would suggest that you *not* enable stalking for hosts and services without thinking things through. Still, its there if you need and want it.

---

# Performance Data

---

## Introduction

Nagios is designed to allow [plugins](#) to return optional performance data in addition to normal status data, as well as allow you to pass that performance data to external applications for processing. A description of the different types of performance data, as well as information on how to go about processing that data is described below...

## Types of Performance Data

There are two basic categories of performance data that can be obtained from Nagios:

1. **Check performance data**
2. **Plugin performance data**

*Check performance data* is internal data that relates to the actual execution of a host or service check. This might include things like service check latency (i.e. how "late" was the service check from its scheduled execution time) and the number of seconds a host or service check took to execute. This type of performance data is available for all checks that are performed. The `$HOSTEXECUTIONTIME$` and `$SERVICEEXECUTIONTIME$` [macros](#) can be used to determine the number of seconds a host or service check was running and the `$HOSTLATENCY$` and `$SERVICELATENCY$` macros can be used to determine how "late" a regularly-scheduled host or service check was.

*Plugin performance data* is external data specific to the plugin used to perform the host or service check. Plugin-specific data can include things like percent packet loss, free disk space, processor load, number of current users, etc. - basically any type of metric that the plugin is measuring when it executes. Plugin-specific performance data is optional and may not be supported by all plugins. As of this writing, no plugins return performance data, although they mostly likely will in the near future. Plugin-specific performance data (if available) can be obtained by using the `$HOSTPERFDATA$` and `$SERVICEPERFDATA$` [macros](#). See below for more information on how plugins can return performance data to Nagios for inclusion in the `$HOSTPERFDATA$` and `$SERVICEPERFDATA$` macros.

## Performance Data Support For Plugins

Normally plugins return a single line of text that indicates the status of some type of measurable data. For example, the `check_ping` plugin might return a line of text like the following:

```
PING ok - Packet loss = 0%, RTA = 0.80 ms
```

With this type of output, the entire line of text is available in the `$HOSTOUTPUT$` or `$SERVICEOUTPUT$` [macros](#) (depending on whether this plugin was used as a host check or service check).

In order to facilitate the passing of plugin-specific performance data to Nagios, the plugin specification has been expanded. If a plugin wishes to pass performance data back to Nagios, it does so by sending the normal text string that it usually would, followed by a pipe character (`|`), and then a string containing one or more performance data metrics. Let's take the `check_ping` plugin as an example and assume that it has been enhanced to return percent packet loss and average round trip time as performance data metrics. A sample plugin output might look like this:

PING ok - Packet loss = 0%, RTA = 0.80 ms | percent\_packet\_loss=0, rta=0.80

When Nagios seems this format of plugin output it will split the output into two parts: everything before the pipe character is considered to be the "normal" plugin output and everything after the pipe character is considered to be the plugin-specific performance data. The "normal" output gets stored in the \$HOSTOUTPUT\$ or \$SERVICEOUTPUT\$ macro, while the optional performance data gets stored in the \$HOSTPERFDATA\$ or \$SERVICEPERFDATA\$ macro. In the example above, the \$HOSTOUTPUT\$ or \$SERVICEOUTPUT\$ macro would contain "PING ok - Packet loss = 0%, RTA = 0.80 ms" (without quotes) and the \$HOSTPERFDATA\$ or \$SERVICEPERFDATA\$ macro would contain "percent\_packet\_loss=0, rta=0.80" (without quotes).

### **Format of Performance Data Output**

The Nagios daemon doesn't directly process performance data, so it doesn't really care what the performance data looks like. There aren't really any inherent limitations on the format or content of the performance data. However, if you are using an external addon to process the performance data (i.e. PerfParse), the addon may be expecting that the plugin returns performance data in a specific format. Check the documentation that comes with the addon for more information. Also, make sure to check the plugin developer guidelines at SourceForge (<http://nagiosplug.sourceforge.net/>) for information on writing plugins.

### **Enabling Performance Data Processing**

If you want to process the performance data that is available from Nagios and the plugins, you'll need to do the following:

1. Enable the [process\\_performance\\_data](#) option.
2. Configure Nagios so that performance data is written to files and/or processed by executing commands.

### **Writing Performance Data To Files**

You can have Nagios write all host and service performance data to files using the [host\\_perfdata\\_file](#) and [service\\_perfdata\\_file](#) options. You can control how the data is written to those files using the [host\\_perfdata\\_file\\_template](#) and [service\\_perfdata\\_file\\_template](#) options. Additionally, you can have Nagios periodically execute commands to process the performance data files using the [host\\_perfdata\\_file\\_processing\\_command](#) and [service\\_perfdata\\_file\\_processing\\_command](#) options.

### **Processing Performance Data Using Commands**

You can have Nagios process host and service performance data by executing commands by using the [host\\_perfdata\\_command](#) or [service\\_perfdata\\_command](#) options. An example command definition that simply writes service performance data to a file is shown below:

```
define command{
  command_name    process-service-perfdata
  command_line    /bin/echo -e "$LASTSERVICECHECKS\t$HOSTNAME\t$SERVICEDESC\t$SERVICESTATES\t$SERVICEATTEMPTS\t$SERVICESTATETYPES\t$SERVICEEXECUTIONTIMES\t$SERVICEDELAY\t$SERVICEOUTPUTS\t$S...
```

---

# Scheduled Downtime

---

## Introduction

Nagios allows you to schedule periods of planned downtime for hosts and service that you're monitoring. This is useful in the event that you actually know you're going to be taking a server down for an upgrade, etc. When a host or service is in a period of scheduled downtime, notifications for that host or service will be suppressed.

## Downtime File

Scheduled host and service downtime is stored in the file you specify by the [downtime\\_file](#) directive in your main configuration file.

## Downtime Retention

Scheduled host and service downtime is automatically preserved across program restarts. When Nagios starts up, it will scan the [downtime file](#), delete any old or invalid entries, and schedule downtime for all valid host and service entries.

## Scheduling Downtime

You can schedule downtime for hosts and service through the [extinfo CGI](#) (either when viewing host or service information). Click in the "Schedule downtime for this host/service" link to actually schedule the downtime.

Once you schedule downtime for a host or service, Nagios will add a comment to that host/service indicating that it is scheduled for downtime during the period of time you indicated. When that period of downtime passes, Nagios will automatically delete the comment that it added. Nice, huh?

## Fixed vs. Flexible Downtime

When you schedule downtime for a host or service through the web interface you'll be asked if the downtime is fixed or flexible. Here's an explanation of how "fixed" and "flexible" downtime differs:

"Fixed" downtime starts and stops at the exact start and end times that you specify when you schedule it. Okay, that was easy enough...

"Flexible" downtime is intended for times when you know that a host or service is going to be down for X minutes (or hours), but you don't know exactly when that'll start. When you schedule flexible downtime, Nagios will start the scheduled downtime sometime between the start and end times you specified. The downtime will last for as long as the duration you specified when you scheduled the downtime. This assumes that the host or service for which you scheduled flexible downtime either goes down (or becomes unreachable) or goes into a non-OK state sometime between the start and end times you specified. The time at which a host or service transitions to a problem state determines the time at which Nagios actually starts the downtime. The downtime will then last for the duration you specified, even if the host or service recovers before the downtime expires. This is done for a very good reason. As we all know, you can think you've got a problem fixed (and restart a server) ten times before it actually works right. Smart, eh?

## Triggered Downtime

When scheduling host or service downtime you have the option of making it "triggered" downtime. What is triggered downtime, you ask? With triggered downtime the start of the downtime is triggered by the start of some other scheduled host or service downtime. This is extremely useful if you're scheduling downtime for a large number of hosts or services and the start time of the downtime period depends on the start time of another downtime entry. For instance, if you schedule flexible downtime for a particular host (because its going down for maintenance), you might want to schedule triggered downtime for all of that host's "children".

### **How Scheduled Downtime Affects Notifications**

When a host or service is in a period of scheduled downtime, Nagios will not allow notifications to be sent out for the host or service. suppression of notifications is accomplished by adding an additional filter to the [notification logic](#). You will *not* see an icon in the CGIs indicating that notifications for that host/service are disabled. When the scheduled downtime has passed, Nagios will allow notifications to be sent out for the host or service as it normally would.

### **Overlapping Scheduled Downtime**

I like to refer to this as the "Oh crap, its not working" syndrome. You know what I'm talking about. You take a server down to perform a "routine" hardware upgrade, only to later realize that the OS drivers aren't working, the RAID array blew up, or the drive imaging failed and left your original disks useless to the world. Moral of the story is that any routine work on a server is quite likely to take three or four times as long as you had originally planned...

Let's take the following scenario:

1. You schedule downtime for host A from 7:30pm-9:30pm on a Monday
2. You bring the server down about 7:45pm Monday evening to start a hard drive upgrade
3. After wasting an hour and a half battling with SCSI errors and driver incompatibilities, you finally get the machine to boot up
4. At 9:15 you realize that one of your partitions is either hosed or doesn't seem to exist anywhere on the drive
5. Knowing you're in for a long night, you go back and schedule additional downtime for host A from 9:20pm Monday evening to 1:30am Tuesday Morning.

If you schedule overlapping periods of downtime for a host or service (in this case the periods were 7:40pm-9:30pm and 9:20pm-1:30am), Nagios will wait until the last period of scheduled downtime is over before it allows notifications to be sent out for that host or service. In this example notifications would be suppressed for host A until 1:30am Tuesday morning.

---



# Using The Embedded Perl Interpreter

---

## Introduction

Stephen Davies has contributed code that allows you to compile Nagios with an embedded Perl interpreter. This may be of interest to you if you rely heavily on plugins written in Perl.

Stanley Hopcroft has worked with the embedded Perl interpreter quite a bit and has commented on the advantages/disadvantages of using it. He has also given several helpful hints on creating Perl plugins that work properly with the embedded interpreter. The majority of this documentation comes from his comments.

It should be noted that "ePN", as used in this documentation, refers to embedded Perl Nagios, or if you prefer, Nagios compiled with an embedded Perl interpreter.

## Advantages

Some advantages of ePN (embedded Perl Nagios) include:

- Nagios will spend much less time running your Perl plugins because it no longer forks to execute the plugin (each time loading the Perl interpreter). Instead, it executes your plugin by making a library call.
- It greatly reduces the system impact of Perl plugins and/or allows you to run more checks with Perl plugin than you otherwise would be able to. In other words, you have less incentive to write plugins in other languages such as C/C++, or Expect/TCL, that are generally recognised to have development times at least an order of magnitude slower than Perl (although they do run about ten times faster also - TCL being an exception).
- If you are not a C programmer, then you can still get a huge amount of mileage out of Nagios by letting Perl do all the heavy lifting without having Nagios slow right down. Note however, that the ePN will not speed up your plugin (apart from eliminating the interpreter load time). If you want fast plugins then consider Perl XSUBs (XS), or C *after* you are sure that your Perl is tuned and that you have a suitable algorithm (Benchmark.pm is *invaluable* for comparing the performance of Perl language elements).
- Using the ePN is an excellent opportunity to learn more about Perl.

## Disadvantages

The disadvantages of ePN (embedded Perl Nagios) are much the same as Apache mod\_perl (i.e. Apache with an embedded interpreter) compared to a plain Apache:

- A Perl program that works *fine* with plain Nagios may *not* work with the ePN. You may have to modify your plugins to get them to work.
- Perl plugins are harder to debug under an ePN than under a plain Nagios.
- Your ePN will have a larger SIZE (memory footprint) than a plain Nagios.
- Some Perl constructs cannot be used or may behave differently than what you would expect.
- You may have to be aware of 'more than one way to do it' and choose a way that seems less attractive or obvious.
- You will need greater Perl knowledge (but nothing very esoteric or stuff about Perl internals - unless your plugin uses XSUBS).

## Target Audience

- Average Perl developers; those with an appreciation of the languages powerful features without knowledge of internals or an in depth knowledge of those features.
- Those with a utilitarian appreciation rather than a great depth of understanding.
- If you are happy with Perl objects, name management, data structures, and the debugger, that's probably sufficient.

### Things you should do when developing a Perl Plugin (ePN or not)

- Always always generate some output
- Use 'use utils' and import the stuff it exports (\$TIMEOUT %ERRORS &print\_revision &support)
- Have a look at how the standard Perl plugins do their stuff e.g.
  - Always exit with \$ERRORS{CRITICAL}, \$ERRORS{OK}, etc.
  - Use getopt to read command line arguments
  - Manage timeouts
  - Call print\_usage (supplied by you) when there are no command line arguments
  - Use standard switch names (eg H 'host', V 'version')

### Things you must do to develop a Perl plugin for ePN

1. <DATA> can not be used; use here documents instead e.g.

```
my $data = <<DATA;
portmapper 100000
portmap 100000
sunrpc 100000
rpcbind 100000
rstatd 100001
rstat 100001
rup 100001
..
DATA

%prognum = map { my($a, $b) = split; ($a, $b) } split(/\n/, $data) ;
```

2. BEGIN blocks will not work as you expect. May be best to avoid.
3. Ensure that it is squeaky clean at compile time i.e.
  - use strict
  - use perl -w (other switches [T notably] may not help)
  - use perl -c
4. Avoid lexical variables (my) with global scope as a means of passing \_\_variable\_\_ data into subroutines. In fact this is \_\_fatal\_\_ if the subroutine is called by the plugin more than once when the check is run. Such subroutines act as 'closures' that lock the global lexicals first value into subsequent calls of the subroutine. If however, your global is read-only (a complicated structure for example) this is not a problem. What Bekman [recommends you do instead](#), is any of the following:
  - make the subroutine anonymous and call it via a code ref e.g.

| turn this                    | into                                   |
|------------------------------|----------------------------------------|
| my \$x = 1 ;                 | my \$x = 1 ;                           |
| sub a { .. Process \$x ... } | \$a_cr = sub { ... Process \$x ... } ; |
| .                            | .                                      |
| .                            | .                                      |
| a ;                          | &\$a_cr ;                              |
| \$x = 2                      | \$x = 2 ;                              |
| a ;                          | &\$a_cr ;                              |

# anon closures \_\_always\_\_ rebind the current lexical value

- put the global lexical and the subroutine using it in their own package (as an object or a module)
- pass info to subs as references or aliases (`\$lex_var` or `$_[n]`)
- replace lexicals with package globals and exclude them from 'use strict' objections with 'use vars qw(global1 global2 ..)'

5. Be aware of where you can get more information.

Useful information can be had from the usual suspects (the O'Reilly books, plus Damien Conways "Object Oriented Perl") but for the really useful stuff in the right context start at Stas Bekman's `mod_perl` guide at <http://perl.apache.org/guide/>.

This wonderful book sized document has nothing whatsoever about Nagios, but all about writing Perl programs for the embedded Perl interpreter in Apache (ie Doug MacEacherns `mod_perl`).

The `perlembed` manpage is essential for context and encouragement.

On the basis that Lincoln Stein and Doug MacEachern know a thing or two about Perl and embedding Perl, their book 'Writing Apache Modules with Perl and C' is almost certainly worth looking at.

6. Be aware that your plugin may return strange values with an ePN and that this is likely to be caused by the problem in item #4 above

7. Be prepared to debug via:

- having a test ePN and
- adding print statements to your plugin to display variable values to STDERR (can't use STDOUT)
- adding print statements to `p1.pl` to display what ePN thinks your plugin is before it tries to run it (vi)
- running the ePN in foreground mode (probably in conjunction with the former recommendations)
- use the 'Deparse' module on your plugin to see how the parser has optimised it and what the interpreter will actually get. (see 'Constants in Perl' by Sean M. Burke, The Perl Journal, Fall 2001)

`perl -MO::Deparse <your_program>`

8. Be aware of what ePN is transforming your plugin too, and if all else fails try and debug the transformed version.

As you can see below `p1.pl` rewrites your plugin as a subroutine called 'hndlr' in the package named 'Embed::<something\_related\_to\_your\_plugin\_file\_name>'.

Your plugin may be expecting command line arguments in `@ARGV` so `p1.pl` also assigns `@_` to `@ARGV`.

This in turn gets 'eval' ed and if the eval raises an error (any parse error and run error), the plugin gets chucked out.

The following output shows how a test ePN transformed the `check_rpc` plugin before attempting to execute it. Most of the code from the actual plugin is not shown, as we are interested in only the transformations that the ePN has made to the plugin). For clarity, transformations are shown in red:

```

        package main;
        use subs 'CORE::GLOBAL::exit';
        sub CORE::GLOBAL::exit { die "ExitTrap: $_[0]"; }
    (Embed::check_5frpc)"; }
    package Embed::check_5frpc; sub hndlr { shift(@_);

@ARGV=@_;
#!/usr/bin/perl -w
#
# check_rpc plugin for Nagios
#
# usage:
#   check_rpc host service
#
# Check if an rpc service is registered and running
# using rpcinfo - $proto $host $prognum 2>&1 |";
#
# Use these hosts.cfg entries as examples
#
# command[check_nfs]=/some/path/libexec/check_rpc $HOSTADDRESS$ nfs
# service[check_nfs]=NFS;24x7;3;5;5;unix-admin;60;24x7;1;1;1;;check_rpc
#
# initial version: 3 May 2000 by Truongchinh Nguyen and Karl DeBisschop
# current status: $Revision: 1.26.2.2 $
#
# Copyright Notice: GPL
#
... rest of plugin code goes here (it was removed for brevity) ...
}

```

9. Don't use 'use diagnostics' in a plugin run by your production ePN. I think it causes \_\_all\_\_ the Perl plugins to return CRITICAL.
10. Consider using a mini embedded Perl C program to check your plugin. This is not sufficient to guarantee your plugin will perform Ok with an ePN but if the plugin fails this test it will certainly fail with your ePN. [ A sample mini ePN is included in the *contrib/* directory of the Nagios distribution for use in testing Perl plugins. Change to the *contrib/* directory and type 'make mini\_epn' to compile it. It must be executed from the same directory that the p1.pl file resides in (this file is distributed with Nagios). ]

### **Compiling Nagios With The Embedded Perl Interpreter**

Okay, you can breathe again now. So do you *still* want to compile Nagios with the embedded Perl interpreter? ;-)

If you want to compile Nagios with the embedded Perl interpreter you need to rerun the configure script with the addition of the *--enable-embedded-perl* option. If you want the embedded interpreter to cache internally compiled scripts, add the *--with-perlcache* option as well. Example:

```
./configure --enable-embedded-perl --with-perlcache ...other options...
```

Once you've rerun the configure script with the new options, make sure to recompile Nagios. You can check to make sure that Nagios has been compiled with the embedded Perl interpreter by executing it with the *-m* command-line argument. Output from executing the command will look something like this (notice that the embedded perl interpreter is listed in the options section):

```

[nagios@firestore ]# ./nagios -m

Nagios 1.0a0
Copyright (c) 1999-2001 Ethan Galstad (nagios@nagios.org)
Last Modified: 07-03-2001
License: GPL

```

## External Data I/O

-----

Object Data:       DEFAULT  
Status Data:       DEFAULT  
Retention Data:     DEFAULT  
Comment Data:       DEFAULT  
Downtime Data:      DEFAULT  
Performance Data:  DEFAULT

## Options

-----

\* Embedded Perl compiler (With caching)

---

# Adaptive Monitoring

---

## Introduction

Nagios allows you to change certain commands and host and service check attributes during runtime. I'll refer to this feature as "adaptive monitoring". Please note that the adaptive monitoring features found in Nagios will probably not be of much use to 99% of users, but they do allow you to do some neat things.

## What Can Be Changed?

The following service check attributes can be changed during runtime:

- Check command (and command arguments)
- Event handler command (and command arguments)
- Check interval
- Max check attempts

The following host check attributes can be changed during runtime:

- Check command (and command arguments)
- Event handler command (and command arguments)
- Check interval
- Max check attempts

The following global attributes can be changed during runtime:

- Global host event handler command (and command arguments)
- Global service event handler command (and command arguments)

## External Commands For Adaptive Monitoring

In order to change global or host- or service-specific attributes during runtime, you must submit the appropriate [external command](#) to Nagios via the [external command file](#). The table below lists the different attributes that may be changed during runtime, along with the external command to accomplish the job.

**NOTE:** When changing check commands or event handler commands, it is important to note that these commands must have been configured using [command definitions](#) before Nagios was started. Any request to change an check or event event handler command to use a command which has not been defined is ignore. Also of note, you specify command arguments along with the actual command name - just seperate individual arguments from the command name (and from each other) using bang (!) characters. More information on how arguments in command definitions are processed during runtime can be found in the documentation on [macros](#).

| Attribute | External Command | Notes |
|-----------|------------------|-------|
|-----------|------------------|-------|

|                              |                                                   |                                                                                                                       |
|------------------------------|---------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------|
| Service check command        | CHANGE_SVC_CHECK_COMMAND: <i>command_name</i>     | Changes the service's current check command to whatever you specify in the <i>command_name</i> argument.              |
| Service event handler        | CHANGE_SVC_EVENT_HANDLER: <i>command_name</i>     | Changes the service's current event handler command to whatever you specify in the <i>command_name</i> argument.      |
| Service check interval       | CHANGE_NORMAL_SVC_CHECK_INTERVAL: <i>interval</i> | Changes the service's normal check interval to be whatever you specify in the <i>interval</i> argument.               |
| Service check retry interval | CHANGE_RETRY_SVC_CHECK_INTERVAL: <i>interval</i>  | Changes the services' retry check interval to be whatever you specify in the <i>interval</i> argument.                |
| Max service check attempts   | CHANGE_MAX_SVC_CHECK_ATTEMPTS: <i>attempts</i>    | Changes the maximum number of check attempts for the service to whatever you specify in the <i>attempts</i> argument. |
| Host check command           | CHANGE_HOST_CHECK_COMMAND: <i>command_name</i>    | Changes the host's current check command to whatever you specify in the <i>command_name</i> argument.                 |

|                              |                                                       |                                                                                                                                       |
|------------------------------|-------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------|
| Host event handler           | CHANGE_HOST_EVENT_HANDLER: <i>command_name</i>        | Changes the host's current event handler command to whatever you specify in the <i>command_name</i> argument.                         |
| Host check interval          | CHANGE_NORMAL_HOST_CHECK_INTERVAL: <i>interval</i>    | Changes the host's check interval to be whatever you specify in the <i>interval</i> argument.                                         |
| Max host check attempts      | CHANGE_MAX_HOST_CHECK_ATTEMPTS: <i>attempts</i>       | Changes the maximum number of check attempts for the host to whatever you specify in the <i>attempts</i> argument.                    |
| Global host event handler    | CHANGE_GLOBAL_HOST_EVENT_HANDLER; <i>command_name</i> | Changes the current <a href="#">global host event handler command</a> to whatever you specify in the <i>command_name</i> argument.    |
| Global service event handler | CHANGE_GLOBAL_SVC_EVENT_HANDLER; <i>command_name</i>  | Changes the current <a href="#">global service event handler command</a> to whatever you specify in the <i>command_name</i> argument. |



# Object Inheritance

---

## Introduction

This documentation attempts to explain object inheritance and how it can be used in [template-based object definitions](#).

One of my primary motivations for adding support for template-based object data was its ability to easily allow object definitions to inherit various properties from other object definitions. Object property inheritance is accomplished through recursion when Nagios processes your configuration files.

If you are still confused about how recursion and inheritance work after reading this, take a look at the sample object config files provided in the distribution. If that still doesn't help, drop an email message with a *detailed* description of your problem to the *nagios-users* mailing list.

## Basics

There are three variables affecting recursion and inheritance that are present in all object definitions. They are indicated in red as follows...

```
define someobjecttype{
    object-specific variables ...
    name                template_name
    use                  name_of_template_to_use
    register             [0/1]
}
```

The first variable is *name*. Its just a "template" name that can be referenced in other object definitions so they can inherit the objects properties/variables. Template names must be unique amongst objects of the same type, so you can't have two or more host definitions that have "hosttemplate" as their template name.

The second variable is *use*. This is where you specify the name of the template object that you want to inherit properties/variables from. The name you specify for this variable must be defined as another object's template named (using the *name* variable).

The third variable is *register*. This variable is used to indicate whether or not the object definition should be "registered" with Nagios. By default, all object definitions are registered. If you are using a partial object definition as a template, you would want to prevent it from being registered (an example of this is provided later). Values are as follows: 0 = do NOT register object definition, 1 = register object definition (this is the default). This variable is NOT inherited; every (partial) object definition used as a template must explicitly set the *register* directive to be 0. This prevents the need to override an inherited *register* directive with a value of 1 for every object that should be registered.

## Local Variables vs. Inherited Variables

One important thing to understand with inheritance is that "local" object variables always take precedence over variables defined in the template object. Take a look at the following example of two host definitions (not all required variables have been supplied):

```

define host{
    host_name          bighost1
    check_command       check-host-alive
    notification_options d,u,r
    max_check_attempts 5
    name               hosttemplate1
}

define host{
    host_name          bighost2
    max_check_attempts 3
    use                hosttemplate1
}

```

You'll note that the definition for host *bighost1* has been defined as having *hosttemplate1* as its template name. The definition for host *bighost2* is using the definition of *bighost1* as its template object. Once Nagios processes this data, the resulting definition of host *bighost2* would be equivalent to this definition:

```

define host{
    host_name          bighost2
    check_command       check-host-alive
    notification_options d,u,r
    max_check_attempts 3
}

```

You can see that the *check\_command* and *notification\_options* variables were inherited from the template object (where host *bighost1* was defined). However, the *host\_name* and *max\_check\_attempts* variables were not inherited from the template object because they were defined locally. Remember, locally defined variables override variables that would normally be inherited from a template object. That should be a fairly easy concept to understand.

## inheritance Chaining

Objects can inherit properties/variables from multiple levels of template objects. Take the following example:

```

define host{
    host_name          bighost1
    check_command       check-host-alive
    notification_options d,u,r
    max_check_attempts 5
    name               hosttemplate1
}

define host{
    host_name          bighost2
    max_check_attempts 3
    use                hosttemplate1
    name               hosttemplate2
}

define host{
    host_name          bighost3
    use                hosttemplate2
}

```

You'll notice that the definition of host *bighost3* inherits variables from the definition of host *bighost2*, which in turn inherits variables from the definition of host *bighost1*. Once Nagios processes this configuration data, the resulting host definitions are equivalent to the following:

```

define host{
    host_name          bighost1
    check_command       check-host-alive
    notification_options d,u,r
    max_check_attempts 5
}

define host{
    host_name          bighost2
    check_command       check-host-alive
    notification_options d,u,r
    max_check_attempts 3
}

define host{
    host_name          bighost3
    check_command       check-host-alive
    notification_options d,u,r
    max_check_attempts 3
}

```

There is no inherent limit on how "deep" inheritance can go, but you'll probably want to limit yourself to at most a few levels in order to maintain sanity.

### Using Incomplete Object Definitions as Templates

It is possible to use incomplete object definitions as templates for use by other object definitions. By "incomplete" definition, I mean that all required variables in the object have not been supplied in the object definition. It may sound odd to use incomplete definitions as templates, but it is in fact recommended that you use them. Why? Well, they can serve as a set of defaults for use in all other object definitions. Take the following example:

```

define host{
    check_command       check-host-alive
    notification_options d,u,r
    max_check_attempts 5
    name                generichosttemplate
    register            0
}

define host{
    host_name          bighost1
    address            192.168.1.3
    use                generichosthosttemplate
}

define host{
    host_name          bighost2
    address            192.168.1.4
    use                generichosthosttemplate
}

```

Notice that the first host definition is incomplete because it is missing the required *host\_name* variable. We don't need to supply a host name because we just want to use this definition as a generic host template. In order to prevent this definition from being registered with Nagios as a normal host, we set the *register* variable to 0.

The definitions of hosts *bighost1* and *bighost2* inherit their values from the generic host definition. The only variable we've chosen to override is the *address* variable. This means that both hosts will have the exact same properties, except for their *host\_name* and *address* variables. Once Nagios processes the config data in the example, the resulting host definitions would be equivalent to specifying the following:

```
define host{
    host_name      bighost1
    address        192.168.1.3
    check_command  check-host-alive
    notification_options d,u,r
    max_check_attempts 5
}

define host{
    host_name      bighost2
    address        192.168.1.4
    check_command  check-host-alive
    notification_options d,u,r
    max_check_attempts 5
}
```

At the very least, using a template definition for default variables will save you a lot of typing. It'll also save you a lot of headaches later if you want to change the default values of variables for a large number of hosts.

---

# Time-Saving Tricks For Object Definitions

or...

"How To Preserve Your Sanity"

---

## Introduction

This documentation attempts to explain how you can exploit the (somewhat) hidden features of [template-based object definitions](#) to save your sanity. How so, you ask? Several types of objects allow you to specify multiple host names and/or hostgroup names in definitions, allowing you to "copy" the object definition to multiple hosts or services. I'll cover each type of object that supports these features separately. For starters, the object types which support this time-saving feature are as follows:

- [Services](#)
- [Service escalations](#)
- [Service dependencies](#)
- [Host escalations](#)
- [Host dependencies](#)
- [Hostgroups](#)

Object types that are not listed above (i.e. timeperiods, commands, etc.) do not support the features I'm about to describe.

## Regular Expression Matching

The examples I give below use "standard" matching of object names. If you wish, you can enable regular expression matching for object names by using the [use\\_regexp\\_matching](#) config option. By default, regular expression matching will only be used in object names that contain the \* and ? wildcard characters. If you want regular expression matching to be used on all object names (regardless of whether or not they contain the \* and ? wildcard characters), enable the [use\\_true\\_regexp\\_matching](#) config option.

Regular expressions can be used in any of the fields used in the examples below (host names, hostgroup names, service names, and servicegroup names).

**NOTE:** Be careful when enabling regular expression matching - you may have to change your config file, since some directives that you might not want to be interpreted as a regular expression just might be! Any problems should become evident once you verify your configuration.

## Service Definitions

**Multiple Hosts:** If you want to create identical [services](#) that are assigned to multiple hosts, you can specify multiple hosts in the *host\_name* directive as follows:

```
define service{
    host_name          HOST1,HOST2,HOST3,...,HOSTN
    service_description SOMESERVICE
    other service directives ...
}
```

The definition above would create a service called *SOMESERVICE* on hosts *HOST1* through *HOSTN*. All the instances of the *SOMESERVICE* service would be identical (i.e. have the same check command, max check attempts, notification period, etc.).

**All Hosts In Multiple Hostgroups:** If you want to create identical services that are assigned to all hosts in one or more hostgroups, you can do so by creating a single service definition. How? The *hostgroup\_name* directive allows you to specify the name of one or more hostgroups that the service should be created for:

```
define service{
    hostgroup_name      HOSTGROUP1,HOSTGROUP2,...,HOSTGROUPN
    service_description  SOMESERVICE
    other service directives ...
}
```

The definition above would create a service called *SOMESERVICE* on all hosts that are members of hostgroups *HOSTGROUP1* through *HOSTGROUPN*. All the instances of the *SOMESERVICE* service would be identical (i.e. have the same check command, max check attempts, notification period, etc.).

**All Hosts:** If you want to create identical services that are assigned to all hosts that are defined in your configuration files, you can use a wildcard in the *host\_name* directive as follows:

```
define service{
    host_name           *
    service_description  SOMESERVICE
    other service directives ...
}
```

The definition above would create a service called *SOMESERVICE* on **all hosts** that are defined in your configuration files. All the instances of the *SOMESERVICE* service would be identical (i.e. have the same check command, max check attempts, notification period, etc.).

### Service Escalation Definitions

**Multiple Hosts:** If you want to create [service escalations](#) for services of the same name/description that are assigned to multiple hosts, you can specify multiple hosts in the *host\_name* directive as follows:

```
define serviceescalation{
    host_name      HOST1,HOST2,HOST3,...,HOSTN
    service_description  SOMESERVICE
    other escalation directives ...
}
```

The definition above would create a service escalation for services called *SOMESERVICE* on hosts *HOST1* through *HOSTN*. All the instances of the service escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

**All Hosts In Multiple Hostgroups:** If you want to create service escalations for services of the same name/description that are assigned to all hosts in one or more hostgroups, you can do use the *hostgroup\_name* directive as follows:

```
define serviceescalation{
    hostgroup_name      HOSTGROUP1,HOSTGROUP2,...,HOSTGROUPN
    service_description  SOMESERVICE
    other escalation directives ...
}
```

The definition above would create a service escalation for services called *SOMESERVICE* on all hosts that are members of hostgroups *HOSTGROUP1* through *HOSTGROUPN*. All the instances of the service escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

**All Hosts:** If you want to create identical service escalations for services of the same name/description that are assigned to all hosts that are defined in your configuration files, you can use a wildcard in the *host\_name* directive as follows:

```
define serviceescalation{
    host_name          *
    service_description SOMESERVICE
    other escalation directives ...
}
```

The definition above would create a service escalation for all services called *SOMESERVICE* on **all hosts** that are defined in your configuration files. All the instances of the service escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

**All Services On Same Host:** If you want to create [service escalations](#) for all services assigned to a particular host, you can use a wildcard in the *service\_description* directive as follows:

```
define serviceescalation{
    host_name          HOST1
    service_description *
    other escalation directives ...
}
```

The definition above would create a service escalation for **all** services on host *HOST1*. All the instances of the service escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

If you feel like being particularly adventurous, you can specify a wildcard in both the *host\_name* and *service\_description* directives. Doing so would create a service escalation for **all services** that you've defined in your configuration files.

**Multiple Services On Same Host:** If you want to create [service escalations](#) for all multiple services assigned to a particular host, you can use a specify more than one service description in the *service\_description* directive as follows:

```
define serviceescalation{
    host_name          HOST1
    service_description SERVICE1, SERVICE2, ..., SERVICEN
    other escalation directives ...
}
```

The definition above would create a service escalation for services *SERVICE1* through *SERVICEN* on host *HOST1*. All the instances of the service escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

**All Services In Multiple Servicegroups:** If you want to create service escalations for all services that belong in one or more servicegroups, you can do use the *servicegroup\_name* directive as follows:

```
define serviceescalation{
    servicegroup_name  SERVICEGROUP1, SERVICEGROUP2, ..., SERVICEGROUPN
    other escalation directives ...
}
```

The definition above would create service escalations for all services that are members of servicegroups *SERVICEGROUP1* through *SERVICEGROUPN*. All the instances of the service escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

## Service Dependency Definitions

**Multiple Hosts:** If you want to create [service dependencies](#) for services of the same name/description that are assigned to multiple hosts, you can specify multiple hosts in the *host\_name* and or *dependent\_host\_name* directives as follows:

```
define servicedependency{
    host_name                HOST1,HOST2
    service_description      SERVICE1
    dependent_host_name      HOST3,HOST4
    dependent_service_description SERVICE2
    other dependency directives ...
}
```

In the example above, service *SERVICE2* on hosts *HOST3* and *HOST4* would be dependent on service *SERVICE1* on hosts *HOST1* and *HOST2*. All the instances of the service dependencies would be identical except for the host names (i.e. have the same notification failure criteria, etc.).

**All Hosts In Multiple Hostgroups:** If you want to create service dependencies for services of the same name/description that are assigned to all hosts in in one or more hostgroups, you can do use the *hostgroup\_name* and/or *dependent\_hostgroup\_name* directives as follows:

```
define servicedependency{
    hostgroup_name           HOSTGROUP1, HOSTGROUP2
    service_description      SERVICE1
    dependent_hostgroup_name HOSTGROUP3, HOSTGROUP4
    dependent_service_description SERVICE2
    other dependency directives ...
}
```

In the example above, service *SERVICE2* on all hosts in hostgroups *HOSTGROUP3* and *HOSTGROUP4* would be dependent on service *SERVICE1* on all hosts in hostgroups *HOSTGROUP1* and *HOSTGROUP2*. Assuming there were five hosts in each of the hostgroups, this definition would be equivalent to creating 100 single service dependency definitions! All the instances of the service dependency would be identical except for the host names (i.e. have the same notification failure criteria, etc.).

**All Services On Same Host:** If you want to create service dependencies for all services assigned to a particular host, you can use a wildcard in the *service\_description* and/or *dependent\_service\_description* directives as follows:

```
define servicedependency{
    host_name                HOST1
    service_description      *
    dependent_host_name      HOST2
    dependent_service_description *
    other dependency directives ...
}
```

In the example above, **all services** on host *HOST2* would be dependent on **all services** on host *HOST1*. All the instances of the service dependencies would be identical (i.e. have the same notification failure criteria, etc.).

**Multiple Services On Same Host:** If you want to create service dependencies for multiple services assigned to a particular host, you can specify more than one service description in the *service\_description* and/or *dependent\_service\_description* directives as follows:



```

define servicedependency{
    host_name                HOST1
    service_description      SERVICE1, SERVICE2, ..., SERVICEN
    dependent_host_name      HOST2
    dependent_service_description SERVICE1, SERVICE2, ..., SERVICEN
    other dependency directives ...
}

```

**All Services In Multiple Servicegroups:** If you want to create service dependencies for all services that belong in one or more servicegroups, you can do use the *servicegroup\_name* and/or *dependent\_servicegroup\_name* directive as follows:

```

define servicedependency{
    servicegroup_name        SERVICEGROUP1, SERVICEGROUP2, ..., SERVICEGROUPN
    dependent_servicegroup_name SERVICEGROUP3, SERVICEGROUP4, ...SERVICEGROUPN
    other escalation directives ...
}

```

## Host Escalation Definitions

**Multiple Hosts:** If you want to create [host escalations](#) for multiple hosts, you can specify multiple hosts in the *host\_name* directive as follows:

```

define hostescalation{
    host_name                HOST1, HOST2, HOST3, ..., HOSTN
    other escalation directives ...
}

```

The definition above would create a host escalation for hosts *HOST1* through *HOSTN*. All the instances of the host escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

**All Hosts In Multiple Hostgroups:** If you want to create host escalations for all hosts in in one or more hostgroups, you can do use the *hostgroup\_name* directive as follows:

```

define hostescalation{
    hostgroup_name          HOSTGROUP1, HOSTGROUP2, ..., HOSTGROUPN
    other escalation directives ...
}

```

The definition above would create a host escalation on all hosts that are members of hostgroups *HOSTGROUP1* through *HOSTGROUPN*. All the instances of the host escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

**All Hosts:** If you want to create identical host escalations for all hosts that are defined in your configuration files, you can use a wildcard in the *host\_name* directive as follows:

```

define hostescalation{
    host_name                *
    other escalation directives ...
}

```

The definition above would create a hosts escalation for **all hosts** that are defined in your configuration files. All the instances of the host escalation would be identical (i.e. have the same contact groups, notification interval, etc.).

## Host Dependency Definitions

**Multiple Hosts:** If you want to create [host dependencies](#) for multiple hosts, you can specify multiple hosts in the *host\_name* and/or *dependent\_host\_name* directives as follows:

```
define hostdependency{
    host_name          HOST1,HOST2
    dependent_host_name HOST3,HOST4,HOST5
    other dependency directives ...
}
```

The definition above would be equivalent to creating six separate host dependencies. In the example above, hosts *HOST3*, *HOST4* and *HOST5* would be dependent upon both *HOST1* and *HOST2*. All the instances of the host dependencies would be identical except for the host names (i.e. have the same notification failure criteria, etc.).

**All Hosts In Multiple Hostgroups:** If you want to create host escalations for all hosts in one or more hostgroups, you can do use the *hostgroup\_name* and /or *dependent\_hostgroup\_name* directives as follows:

```
define hostdependency{
    hostgroup_name          HOSTGROUP1, HOSTGROUP2
    dependent_hostgroup_name HOSTGROUP3, HOSTGROUP4
    other dependency directives ...
}
```

In the example above, all hosts in hostgroups *HOSTGROUP3* and *HOSTGROUP4* would be dependent on all hosts in hostgroups *HOSTGROUP1* and *HOSTGROUP2*. All the instances of the host dependencies would be identical except for host names (i.e. have the same notification failure criteria, etc.).

## Hostgroups

**All Hosts:** If you want to create a hostgroup that has all hosts that are defined in your configuration files as members, you can use a wildcard in the *members* directive as follows:

```
define hostgroup{
    hostgroup_name          HOSTGROUP1
    members                  *
    other hostgroup directives ...
}
```

The definition above would create a hostgroup called *HOSTGROUP1* that has all **all hosts** that are defined in your configuration files as members.

---

# UCD-SNMP (NET-SNMP) Integration

---

**Note:** Nagios is not designed to be a replacement for a full-blown SNMP management application like HP OpenView or [OpenNMS](#). However, you can set things up so that SNMP traps received by a host on your network can generate alerts in Nagios. Here's how...

## Introduction

This example explains how to easily generate alerts in Nagios for SNMP traps that are received by the [UCD-SNMP](#) *snmptrapd* daemon. These directions assume that the host which is receiving SNMP traps is not the same host on which Nagios is running. If your monitoring box is the same box that is receiving SNMP traps you will need to make a few modifications to the examples I provide. Also, I am assuming that you have installed the [nsca daemon](#) on your monitoring server and the nsca client (*send\_nsca*) on the machine that is receiving SNMP traps.

For the purposes of this example, I will be describing how I setup Nagios to generate alerts from SNMP traps received by the ArcServe backup jobs running on my Novell servers. I wanted to get notified when backups failed, so this worked very nicely for me. You'll have to tweak the examples in order to make it suit your needs.

## Additional Software

Translating SNMP traps into Nagios events can be a bit tedious. If you'd like to make it easier, you might want to check out Alex Burger's SNMP Trap Translator project located at <http://www.snmpptt.org> which, combined with Net-SNMP, provides a more enhanced trap handling system. The snmpptt documentation includes integration details for Nagios.

## Defining The Service

First off you're going to have to define a service in your [object configuration file](#) for the SNMP traps (in this example, I am defining a service for ArcServe backup jobs). Assuming that the host that the alerts are originating from is called **novellserver**, a sample service definition might look something like this:

```
define service{
    host_name                novellserver
    service_description      ArcServe Backup
    is_volatile              1
    active_checks_enabled    0
    passive_checks_enabled   1
    max_check_attempts       1
    contact_groups            novell-backup-admins
    notification_interval    120
    notification_period       24x7
    notification_options      w,u,c,r
    check_command             check_none
}
```

Important things to note are the fact that this service has the *volatile* option enabled. We want this option enabled because we want a notification to be generated for every alert that comes in. Also of note is the fact that active checks are disabled for the service, while passive checks are enabled. This means that the service will never be actively checked - all alert information will have to be sent in passively by the *nsca client* on the SNMP management host (in my example, it will be called **firestorm**).

## ArcServe and Novell SNMP Configuration

In order to get ArcServe (and my Novell server) to send SNMP traps to my management host, I had to do the following:

1. Modify the ArcServe autopilot job to send SNMP traps on job failures, successes, etc.
2. Edit `SYS:\ETC\TRAPTARG.CFG` and add the IP address of my management host (the one receiving the SNMP traps)
3. Load `SNMP.NLM`
4. Load `ALERT.NLM` to facilitate the actual sending of the SNMP traps

### SNMP Management Host Configuration

On my Linux SNMP management host (**firestorm**), I installed the [UCD-SNMP](#) (NET-SNMP) software. Once the software was installed I had to do the following:

1. Install the ArcServe MIBs (included on the ArcServe installation CD)
2. Edit the `snmptrapd` configuration file (`/etc/snmp/snmptrapd.conf`) to define a trap handler for ArcServe alerts. This is detailed below.
3. Start the `snmptrapd` daemon to listen for incoming SNMP traps

In order to have the `snmptrapd` daemon route ArcServe SNMP traps to our Nagios host, we've got to define a traphandler in the `/etc/snmp/snmptrapd.conf` file. In my setup, the config file looked something like this:

```
#####
# ArcServe SNMP Traps
#####

# Tape format failures
traphandle ARCserve-Alarm-MIB::arcServetrap9 /usr/local/nagios/libexec/eventhandlers/handle-arcserve-trap 9

# Failure to read tape header
traphandle ARCserve-Alarm-MIB::arcServetrap10 /usr/local/nagios/libexec/eventhandlers/handle-arcserve-trap 10

# Failure to position tape
traphandle ARCserve-Alarm-MIB::arcServetrap11 /usr/local/nagios/libexec/eventhandlers/handle-arcserve-trap 11

# Cancelled jobs
traphandle ARCserve-Alarm-MIB::arcServetrap12 /usr/local/nagios/libexec/eventhandlers/handle-arcserve-trap 12

# Successful jobs
traphandle ARCserve-Alarm-MIB::arcServetrap13 /usr/local/nagios/libexec/eventhandlers/handle-arcserve-trap 13

# Imcomplete jobs
traphandle ARCserve-Alarm-MIB::arcServetrap14 /usr/local/nagios/libexec/eventhandlers/handle-arcserve-trap 14

# Job failures
traphandle ARCserve-Alarm-MIB::arcServetrap15 /usr/local/nagios/libexec/eventhandlers/handle-arcserve-trap 15
```

This example assumes that you have a `/usr/local/nagios/libexec/eventhandlers/` directory on your SNMP mangement host and that the `handle-arcserve-trap` script exists there. You can modify these to fit your setup. Anyway, the `handle-arcserve-trap` script on my management host looked something like this:

```
#!/bin/sh

# Arguments:
# $1 = trap type

# First line passed from snmptrapd is FQDN of host that sent the trap
read host

# Given a FQDN, get the short name of the host as it is setup in Nagios
hostname="unknown"
case $host in
    novellserver.mylocaldomain.com)
        hostname="novellserver"
        ;;
```

```

nt.mylocaldomain.com)
    hostname="ntserver"
    ;;
esac

# Get severity level (OK, WARNING, UNKNOWN, or CRITICAL) and plugin output based on trape type
state=-1
output="No output"
case "$1" in

    # failed to format tape - critical
    11)
        output="Critical: Failed to format tape"
        state=2
        ;;

    # failed to read tape header - critical
    10)
        output="Critical: Failed to read tape header"
        state=2
        ;;

    # failed to position tape - critical
    11)
        output="Critical: Failed to position tape"
        state=2
        ;;

    # backup cancelled - warning
    12)
        output="Warning: ArcServe backup operation cancelled"
        state=1
        ;;

    # backup success - ok
    13)
        output="Ok: ArcServe backup operation successful"
        state=0
        ;;

    # backup incomplete - warning
    14)
        output="Warning: ArcServe backup operation incomplete"
        state=1
        ;;

    # backup failure - critical
    15)
        output="Critical: ArcServe backup operation failed"
        state=2
        ;;
esac

# Submit passive check result to monitoring host
/usr/local/nagios/libexec/eventhandlers/submit_check_result $hostname "ArcServe Backup" $state "$output"

exit 0

```

Notice that the *handle-arcserve-trap* script calls the *submit\_check\_result* script to actually send the alert back to the monitoring host. Assuming your monitoring host is called **monitor**, the *submit\_check\_result* script might look like this (you'll have to modify this to specify the proper location of the *send\_nsca* program on your management host):

```

#!/bin/sh

# Arguments
#   $1 = name of host in service definition
#   $2 = name/description of service in service definition
#   $3 = return code
#   $4 = output

/bin/echo -e "$1\t$2\t$3\t$4\n" | /usr/local/nagios/bin/send_nsca monitor -c /usr/local/nagios/etc/send_nsca.cfg

```

## **Finishing Up**

You've now configured everything you need to, so all you have to do is restart the Nagios on your monitoring server. That's it! You should be getting alerts in Nagios whenever ArcServe jobs fail, succeed, etc.

---

# TCP Wrapper Integration

---

## Introduction

This example explains how to easily generate alerts in Nagios for connection attempts that are rejected by TCP wrappers. These directions assume that the host which you are generating alerts for (i.e. the host you are using TCP wrappers on) is not the same host on which Nagios is running. If you want to generate alerts on the same host that Nagios is running you will need to make a few modifications to the examples I provide. Also, I am assuming that you have installed the [nsca daemon](#) on your monitoring server and the nsca client (*send\_nsca*) on the machine that you are generating TCP wrapper alerts from.

## Defining The Service

First off you're going to have to define a service in your [object configuration file](#) for the TCP wrapper alerts. Assuming that the host that the alerts are originating from is called **firestorm**, a sample service definition might look something like this:

```
define service{
    host_name                firestorm
    service_description      TCP Wrappers
    is_volatile              1
    active_checks_enabled    0
    passive_checks_enabled   1
    max_check_attempts       1
    contact_groups           security-admins
    notification_interval    120
    notification_period       24x7
    notification_options     w,u,c,r
    check_command             check_none
}
```

Important things to note are the fact that this service has the *volatile* option enabled. We want this option enabled because we want a notification to be generated for every alert that comes in. Also of note is the fact that active checks of the service are disabled, while passive checks are enabled. This means that the service will never be actively checked - all alert information will have to be sent in passively by the *nsca client* on the **firestorm** host.

## Configuring TCP Wrappers

Now you're going to have to modify the */etc/hosts.deny* file on the host called **firestorm**. In order to have the TCP wrappers send an alert to the monitoring host whenever a connection attempt is denied, you'll have to add a line similar to the following:

```
ALL: ALL: RFC931: twist (/usr/local/nagios/libexec/eventhandlers/handle_tcp_wrapper %h %d) &
```

This line assumes that there is a script called *handle\_tcp\_wrapper* in the */usr/local/nagios/libexec/eventhandlers/* directory on **firestorm**. The directory and script name can be changed to whatever you want.

## Writing The Script

The last thing you need to do is write the *handle\_tcp\_wrapper* script on **firestorm** that will send the alert back to the monitoring host. It might look something like this:

```
#!/bin/sh

/usr/local/nagios/libexec/eventhandlers/submit_check_result firestorm "TCP Wrappers" 2 "Denied $2-$1" > /dev/null 2> /dev/null
```

Notice that the *handle\_tcp\_wrapper* script calls the *submit\_check\_result* script to actually send the alert back to the monitoring host. Assuming your monitoring host is called **monitor**, the *submit\_check\_result* script might look like this (you'll have to modify this to specify the proper location of the *send\_nsca* program on **firestorm**):

```
#!/bin/sh

# Arguments
#   $1 = name of host in service definition
#   $2 = name/description of service in service definition
#   $3 = return code
#   $4 = output

/bin/echo -e "$1\t$2\t$3\t$4\n" | /usr/local/nagios/bin/send_nsca monitor -c /usr/local/nagios/etc/send_nsca.cfg
```

## Finishing Up

You've now configured everything you need to, so all you have to do is restart the *inetd* process on **firestorm** and restart Nagios on your monitoring server. That's it! When the TCP wrappers on **firestorm** deny a connection attempt, you should be getting alerts in Nagios. The plugin output for the alert will look something like the following:

```
Denied sshd2-sdn-ar-002mnminnP321.dialsprint.net
```

---



# Securing Nagios

---

## Introduction

This is intended to be a brief overview of some things you should keep in mind when installing Nagios, so as to not set it up in an insecure manner. This document is new, so if anyone has additional notes or comments on securing Nagios, please drop me a note at [nagios@nagios.org](mailto:nagios@nagios.org)

## Do Not Run Nagios As Root!

Nagios doesn't need to run as root, so don't do it. Even if you start Nagios at boot time with an init script, you can force it to drop privileges after startup and run as another user/group by using the [nagios\\_user](#) and [nagios\\_group](#) directives in the main config file.

If you need to execute event handlers or plugins which require root access, you might want to try using [sudo](#).

## Enable External Commands Only If Necessary

By default, [external commands](#) are disabled. This is done to prevent an admin from setting up Nagios and unknowingly leaving its command interface open for use by "others".. If you are planning on using [event handlers](#) or issuing commands from the web interface, you will have to enable external commands. If you aren't planning on using event handlers or the web interface to issue commands, I would recommend leaving external commands disabled.

## Set Proper Permissions On The External Command File

If you enable [external commands](#), make sure you set proper permissions on the `/usr/local/nagios/var/rw` directory. You only want the Nagios user (usually *nagios*) and the web server user (usually *nobody*) to have permissions to write to the command file. If you've installed Nagios on a machine that is dedicated to monitoring and admin tasks and is not used for public accounts, that should be fine.

If you've installed it on a public or multi-user machine, allowing the web server user to have write access to the command file can be a security problem. After all, you don't want just any user on your system controlling Nagios through the external command file. In this case, I would suggest only granting write access on the command file to the *nagios* user and using something like [CGIWrap](#) to run the CGIs as the *nagios* user instead of *nobody*.

Instructions on setting up permissions for the external command file can be found [here](#).

## Require Authentication In The CGIs

I would strongly suggest requiring authentication for accessing the CGIs. Once you do that, read the documentation on the default rights that authenticated contacts have, and only authorize specific contacts for additional rights as necessary. Instructions on setting up authentication and configuring authorization rights can be found [here](#). If you disable the CGI authentication features using the [use\\_authentication](#) directive in the CGI config file, the [command CGI](#) will refuse to write any commands to the [external command file](#). After all, you don't want the world to be able to control Nagios do you?

## Use Full Paths In Command Definitions

When you define commands, make sure you specify the *full path* to any scripts or binaries you're executing.

## **Hide Sensitive Information With \$USERn\$ Macros**

The CGIs read the [main config file](#) and [object config file\(s\)](#), so you don't want to keep any sensitive information (usernames, passwords, etc) in there. If you need to specify a username and/or password in a command definition use a \$USERn\$ [macro](#) to hide it. \$USERn\$ macros are defined in one or more [resource files](#). The CGIs will not attempt to read the contents of resource files, so you can set more restrictive permissions (600 or 660) on them. See the sample *resource.cfg* file in the base of the Nagios distribution for an example of how to define \$USERn\$ macros.

## **Strip Dangerous Characters From Macros**

Use the [illegal\\_macro\\_output\\_chars](#) directive to strip dangerous characters from the \$HOSTOUTPUT\$, \$SERVICEOUTPUT\$, \$HOSTPERFDATA\$, and \$SERVICEPERFDATA\$ macros before they're used in notifications, etc. Dangerous characters can be anything that might be interpreted by the shell, thereby opening a security hole. An example of this is the presence of backtick (`) characters in the \$HOSTOUTPUT\$, \$SERVICEOUTPUT\$, \$HOSTPERFDATA\$, and/or \$SERVICEPERFDATA\$ macros, which could allow an attacker to execute an arbitrary command as the nagios user (one good reason not to run Nagios as the root user).

---

# Tuning Nagios For Maximum Performance

---

## Introduction

So you've finally got Nagios up and running and you want to know how you can tweak it a bit... Here are a few things to look at for optimizing Nagios. Let me know if you think of any others...

### Optimization Tips:

1. **Use aggregated status updates.** Enabling aggregated status updates (with the [aggregate\\_status\\_updates](#) option) will greatly reduce the load on your monitoring host because it won't be constantly trying to update the [status log](#). This is especially recommended if you are monitoring a large number of services. The main trade-off with using aggregated status updates is that changes in the states of hosts and services will not be reflected immediately in the status file. This may or may not be a big concern for you.
2. **Use a ramdisk for holding status data.** If you're using the standard [status log](#) and you're *not* using aggregated status updates, consider putting the directory where the status log is stored on a ramdisk. This will speed things up quite a bit (in both the core program and the CGIs) because it saves a lot of interrupts and disk thrashing.
3. **Check service latencies to determine best value for maximum concurrent checks.** Nagios can restrict the number of maximum concurrently executing service checks to the value you specify with the [max\\_concurrent\\_checks](#) option. This is good because it gives you some control over how much load Nagios will impose on your monitoring host, but it can also slow things down. If you are seeing high latency values (> 10 or 15 seconds) for the majority of your service checks (via the [extinfo CGI](#)), you are probably starving Nagios of the checks it needs. That's not Nagios's fault - it's yours. Under ideal conditions, all service checks would have a latency of 0, meaning they were executed at the exact time that they were scheduled to be executed. However, it is normal for some checks to have small latency values. I would recommend taking the minimum number of maximum concurrent checks reported when running Nagios with the `-s` command line argument and doubling it. Keep increasing it until the average check latency for your services is fairly low. More information on service check scheduling can be found [here](#).
4. **Use passive checks when possible.** The overhead needed to process the results of [passive service checks](#) is much lower than that of "normal" active checks, so make use of that piece of info if you're monitoring a slew of services. It should be noted that passive service checks are only really useful if you have some external application doing some type of monitoring or reporting, so if you're having Nagios do all the work, this won't help things.
5. **Avoid using interpreted plugins.** One thing that will significantly reduce the load on your monitoring host is the use of compiled (C/C++, etc.) plugins rather than interpreted script (Perl, etc) plugins. While Perl scripts and such are easy to write and work well, the fact that they are compiled/interpreted at every execution instance can significantly increase the load on your monitoring host if you have a lot of service checks. If you want to use Perl plugins, consider compiling them into true executables using `perlcc(1)` (a utility which is part of the standard Perl distribution) or compiling Nagios with an embedded Perl interpreter (see below).
6. **Use the embedded Perl interpreter.** If you're using a lot of Perl scripts for service checks, etc., you will probably find that compiling an embedded Perl interpreter into the Nagios binary will speed things up. In order to compile in the embedded Perl interpreter, you'll need to supply the `--enable-embedded-perl` option to the configure script before you compile Nagios. Also, if you use the `--with-perlcache` option, the compiled version of all Perl scripts processed by the embedded interpreter will be cached for later reuse.
7. **Optimize host check commands.** If you're checking host states using the `check_ping` plugin you'll find that host checks will be performed much faster if you break up the checks. Instead of specifying a `max_attempts` value of 1 in the host definition and having the `check_ping` plugin send 10 ICMP

packets to the host, it would be much faster to set the *max\_attempts* value to 10 and only send out 1 ICMP packet each time. This is due to the fact that Nagios can often determine the status of a host after executing the plugin once, so you want to make the first check as fast as possible. This method does have its pitfalls in some situations (i.e. hosts that are slow to respond may be assumed to be down), but I you'll see faster host checks if you use it. Another option would be to use a faster plugin (i.e. *check\_fping*) as the *host\_check\_command* instead of *check\_ping*.

8. **Don't schedule regular host checks.** Do NOT schedule regular checks of hosts unless absolutely necessary. There are not many reasons to do this, as host checks are performed on-demand as needed. To disable regular checks of a host, set the *check\_interval* directive in the [host definition](#) to 0. If you do need to have regularly scheduled host checks, try to use a longer check interval and make sure your host checks are optimized (see above).
  9. **Don't use aggressive host checking.** Unless you're having problems with Nagios recognizing host recoveries, I would recommend *not* enabling the [use\\_aggressive\\_host\\_checking](#) option. With this option turned off host checks will execute much faster, resulting in speedier processing of service check results. However, host recoveries can be missed under certain circumstances when this is turned off. For example, if a host recovers and all of the services associated with that host stay in non-OK states (and don't "wobble" between different non-OK states), Nagios may miss the fact that the host has recovered. A few people may need to enable this option, but the majority don't and I would recommend *not* using it unless you find it necessary...
  10. **Increase external command check interval.** If you're processing a lot of external commands (i.e. passive checks in a [distributed setup](#), you'll probably want to set the [command\\_check\\_interval](#) variable to -1. This will cause Nagios to check for external commands as often as possible. This is important because most systems have small pipe buffer sizes (i.e. 4KB). If Nagios doesn't read the data from the pipe fast enough, applications that write to the external command file (i.e. the [NSCA daemon](#)) will block and wait until there is enough free space in the pipe to write their data.
  11. **Optimize hardware for maximum performance.** Your system configuration and your hardware setup are going to directly affect how your operating system performs, so they'll affect how Nagios performs. The most common hardware optimization you can make is with your hard drives. CPU and memory speed are obviously factors that affect performance, but disk access is going to be your biggest bottleneck. Don't store plugins, the status log, etc on slow drives (i.e. old IDE drives or NFS mounts). If you've got them, use UltraSCSI drives or fast IDE drives. An important note for IDE/Linux users is that many Linux installations do not attempt to optimize disk access. If you don't change the disk access parameters (by using a utility like **hdparam**), you'll lose out on a **lot** of the speedy features of the new IDE drives.
-

# Using The Nagiosstats Utility

---

## Introduction

A utility called *nagiosstats* is included in the Nagios distribution. It is compiled and installed along with the main Nagios daemon.

The *nagiosstats* utility allows you to obtain various information about a running Nagios process. You can obtain information either in human-readable or MRTG-compatible format.

## Usage Information

You can run the *nagiosstats* utility with the **--help** option to get usage information:

```
[nagios@lanman ~]# /usr/local/nagios/bin/nagiosstats --help
```

```
Nagios Stats 2.0a1
Copyright (c) 2003 Ethan Galstad (nagios@nagios.org)
Last Modified: 11-18-2003
License: GPL
```

```
Usage: /usr/local/nagios/bin/nagiosstats [options]
```

### Startup:

```
-V, --version      display program version information and exit.
-L, --license      display license information and exit.
-h, --help         display usage information and exit.
```

### Input file:

```
-c, --config=FILE  specifies location of main Nagios config file.
```

### Output:

```
-m, --mrtg         display output in MRTG compatible format.
-d, --data=VARS    comma-seperated list of variables to output in MRTG
                   (or compatible) format. See possible values below.
                   Percentages are rounded, times are in milliseconds.
```

### MRTG DATA VARIABLES (-d option):

```
NUMSERVICES      total number of services.
NUMHOSTS         total number of services.
NUMSVCOK         number of services OK.
NUMSVCWARN       number of services WARNING.
NUMSVCUNKN       number of services UNKNOWN.
NUMSVCCRIT       number of services CRITICAL.
NUMSVCPROB       number of service problems (WARNING, UNKNOWN or CRITIAL).
NUMHSTUP         number of hosts UP.
NUMHSTDOWN       number of hosts DOWN.
NUMHSTUNR        number of hosts UNREACHABLE.
NUMHSTPROB       number of host problems (DOWN or UNREACHABLE).
xxxACTSVCLAT     MIN/MAX/AVG active service check latency (ms).
xxxACTSVCEXT     MIN/MAX/AVG active service check execution time (ms).
xxxACTSVCPCSC    MIN/MAX/AVG active service check % state change.
xxxPSVSVCPCSC    MIN/MAX/AVG passive service check % state change.
xxxSVCPCSC       MIN/MAX/AVG service check % state change.
xxxACTHSTLAT     MIN/MAX/AVG active host check latency (ms).
xxxACTHSTEXT     MIN/MAX/AVG active host check execution time (ms).
xxxACTHSTPCSC    MIN/MAX/AVG active host check % state change.
xxxPSVHSTPCSC    MIN/MAX/AVG passive host check % state change.
xxxHSTPCSC       MIN/MAX/AVG host check % state change.
NUMACTHSTCHKxM   number of active host checks in last 1/5/15/60 minutes.
NUMPSVHSTCHKxM   number of passive host checks in last 1/5/15/60 minutes.
NUMACTSVCCHKxM   number of active service checks in last 1/5/15/60 minutes.
```

NUMPSVSVCCCHKxM      number of passive service checks in last 1/5/15/60 minutes.

Note: Replace x's in MRTG variable names with 'MIN', 'MAX', 'AVG', or the  
the appropriate number (i.e. '1', '5', '15', or '60').

[nagios@lanman ~]#

## **Human-Readable Output**

For normal operation, run the *nagiosstats* utility, specifying only the config file location as an argument, as follows:

[nagios@lanman ~]# /usr/local/nagios/bin/nagiosstats -c /usr/local/nagios/etc/nagios.cfg

Nagios Stats 2.0a1  
Copyright (c) 2003 Ethan Galstad (nagios@nagios.org)  
Last Modified: 11-18-2003  
License: GPL

### CURRENT STATUS DATA

-----  
Status File:                                /usr/local/nagios/var/status.dat  
Status File Age:                            0d 0h 0m 13s  
Status File Version:                       2.0-very-pre-alpha

Program Running Time:                      14d 17h 19m 13s

Total Services:                           32  
Services Checked:                          32  
Services Scheduled:                        29  
Active Service Checks:                    29  
Passive Service Checks:                   3  
Total Service State Change:               0.000 / 65.530 / 2.930 %  
Active Service Latency:                   0.048 / 14.837 / 1.035 %  
Active Service Execution Time:            0.076 / 60.006 / 4.301 sec  
Active Service State Change:              0.000 / 10.530 / 0.762 %  
Active Services Last 1/5/15/60 min:      1 / 13 / 29 / 29  
Passive Service State Change:             0.000 / 65.530 / 23.883 %  
Passive Services Last 1/5/15/60 min:      0 / 0 / 0 / 0  
Services Ok/Warn/Unk/Crit:               23 / 5 / 1 / 3  
Services Flapping:                        1  
Services In Downtime:                      0

Total Hosts:                              9  
Hosts Checked:                             9  
Hosts Scheduled:                           9  
Active Host Checks:                        9  
Passive Host Checks:                       0  
Total Host State Change:                   0.000 / 28.420 / 4.034 %  
Active Host Latency:                       0.000 / 15.741 / 5.443 %  
Active Host Execution Time:               1.022 / 10.032 / 3.047 sec  
Active Host State Change:                  0.000 / 28.420 / 4.034 %  
Active Hosts Last 1/5/15/60 min:          0 / 8 / 9 / 9  
Passive Host State Change:                  0.000 / 0.000 / 0.000 %  
Passive Hosts Last 1/5/15/60 min:          0 / 0 / 0 / 0  
Hosts Up/Down/Unreach:                   7 / 1 / 1  
Hosts Flapping:                           0  
Hosts In Downtime:                        0

[nagios@lanman ~]#

As you can see, the utility displays a number of different metrics pertaining to the Nagios process. Metrics which have multiple values are (unless otherwise specified) min, max and average values for that particular metric.

## **MRTG Integration**

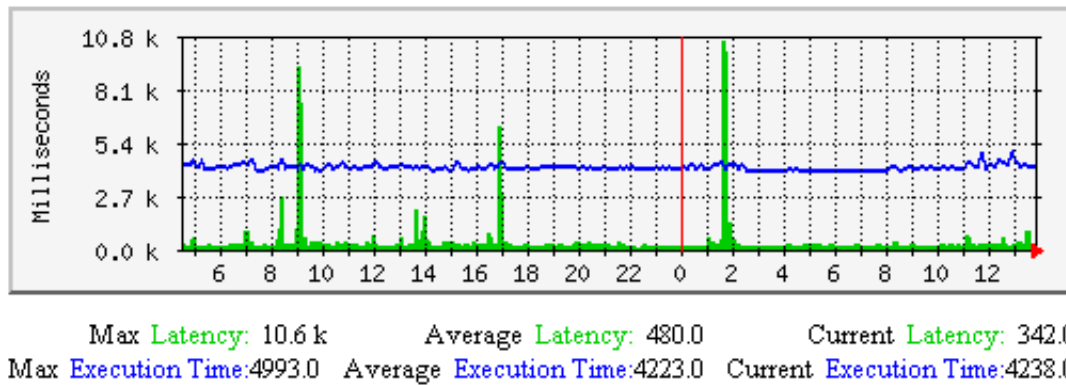
You can use the *nagiosstats* utility to display various Nagios metrics using MRTG (or other compatible program). To do so, run the *nagiosstats* utility using the **--mrtg** and **--data** arguments. The **--data** argument is used to specify what statistics should be graphed. Possible values for the **--data** argument can be found by running the *nagiosstats* utility with the **--help** option.

Here's an MRTG config file snippet for using the *nagiosstats* utility for graphing average service latency and execution time.

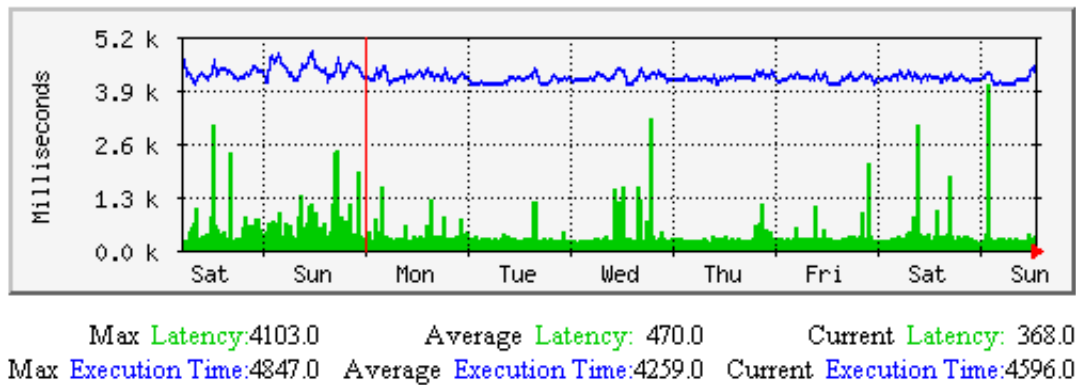
```
# Service Latency and Execution Time
Target[nagios-a]: '/usr/local/nagios/bin/nagiosstats --mrtg --data=AVGACTSVCLAT,AVGACTSVCEXT'
MaxBytes[nagios-a]: 100000
Title[nagios-a]: Average Service Check Latency and Execution Time
PageTop[nagios-a]: <H1>Average Service Check Latency and Execution Time</H1>
Options[nagios-a]: growright,gauge,nopercent
YLegend[nagios-a]: Milliseconds
ShortLegend[nagios-a]: &nbsp;
LegendI[nagios-a]: &nbsp;Latency:
Legend0[nagios-a]: &nbsp;Execution Time:
Legend1[nagios-a]: Latency
Legend2[nagios-a]: Execution Time
Legend3[nagios-a]: Maximal 5 Minute Latency
Legend4[nagios-a]: Maximal 5 Minute Execution Time
```

The MRTG graphs generated from the above config snippet look like this:

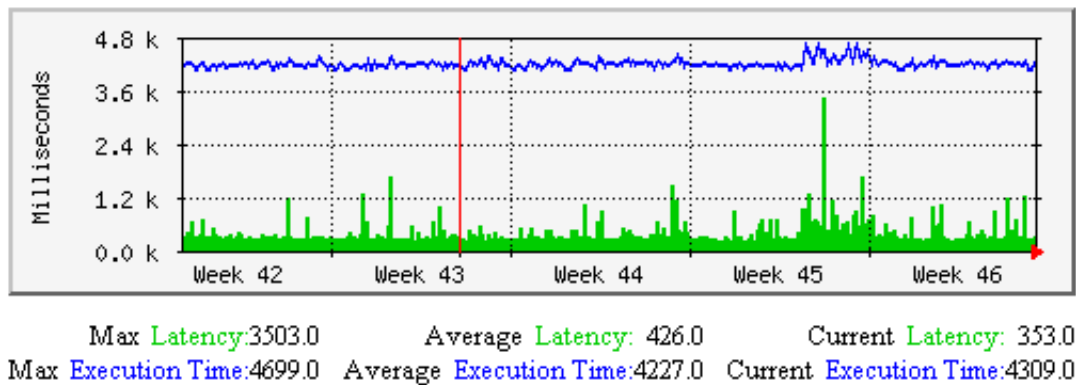
### 'Daily' Graph (5 Minute Average)



### 'Weekly' Graph (30 Minute Average)



### 'Monthly' Graph (2 Hour Average)





# Using Macros In Commands

---

## Macros

One of the features available in Nagios is the ability to use macros in command definitions. Immediately prior to the execution of a command, Nagios will replace all macros in the command with their corresponding values. This allows you to define a few generic commands to handle all your needs.

## Macro Substitution

Before any commands (host and service checks, notifications, event handlers, etc.) are executed, Nagios will replace any macros it finds in the command definition with their corresponding values.

When you use host and service macros in command definitions, they refer to values for the host or service for which the command is being run. Let's try an example. Assuming we are using a host definition and a *check\_ping* command defined like this:

```
define host{
    host_name          linuxbox
    address             192.168.1.2
    check_command       check_ping
    ...
}

define command{
    command_name       check_ping
    command_line        /usr/local/nagios/libexec/check_ping -H $HOSTADDRESS$ -w 100.0,90% -c 200.0,60%
}
```

the expanded/final command line to be executed for the host's check command would look like this:

```
/usr/local/nagios/libexec/check_ping -H 192.168.1.2 -w 100.0,90% -c 200.0,60%
```

You can pass arguments to commands as well, which is quite handy if you'd like to keep your command definitions rather generic. Arguments are specified in the object (i.e. host or service) definition, by separating them from the command name with exclamation points (!) like so:

```
define service{
    host_name          linuxbox
    service_description PING
    ...
    check_command       check_ping!200.0,80%!400.0,40%
    ...
}
```

In the example above, the service check command has two arguments (which can be referenced with [\\$ARGn\\$](#) macros). The `$ARG1$` macro will be "200.0,80%" and `$ARG2$` will be "400.0,40%" (both without quotes). Assuming we are using the host definition given earlier and a *check\_ping* command defined like this:

```
define command{
    command_name       check_ping
    command_line        /usr/local/nagios/libexec/check_ping -H $HOSTADDRESS$ -w $ARG1$ -c $ARG2$
}
```

the expanded/final command line to be executed for the service's check command would look like this:

## **On-Demand Macros**

Normally when you use host and service macros in command definitions, they refer to values for the host or service for which the command is being run. For instance, if a host check command is being executed for a host named "linuxbox", all the host macros listed in the table below will refer to values for that host ("linuxbox").

If you would like to reference values for another host or service in a command (for which the command is not being run), you can use what are called "on-demand" macros. On-demand macros look like normal macros, except for the fact that they contain an identifier for the host or service from which they should get their value. Here's the basic format for on-demand macros:

- `$HOSTMACRO:host_name$`
- `$SERVICEMACRO:host_name:service_description$`

Note that the macro name is separated from the host or service identifier by a colon (:). For on-demand service macros, the service identifier consists of both a host name and a service description - these are separated by a colon (:) as well.

Examples of on-demand host and service macros follow:

```
$HOSTDOWNTIME:myhost$  
$SERVICESTATEID:novellserver:DS Database$
```

## **Macro Cleansing**

Some macros are stripped of potentially dangerous shell metacharacters before being substituted into commands to be executed. Which characters are stripped from the macros depends on the setting of the [illegal\\_macro\\_output\\_chars](#) directive. The following macros are stripped of potentially dangerous characters:

1. `$HOSTOUTPUT$`
2. `$HOSTPERFDATA$`
3. `$HOSTACKAUTHOR$`
4. `$HOSTACKCOMMENT$`
5. `$SERVICEOUTPUT$`
6. `$SERVICEPERFDATA$`
7. `$SERVICEACKAUTHOR$`
8. `$SERVICEACKCOMMENT$`

## **Macros as Environment Variables**

Starting with Nagios 2.0, most macros have been made available as environment variables. This means that scripts that are run from Nagios (i.e. service and host check commands, notification commands, etc.) can reference these macros directly as standard environment variables. For purposes of security and sanity, `$USERn$` and "on-demand" host and service macros are not made available as environment variables. Environment variables that contain macros are named the same as their corresponding macro names (listed below), with "NAGIOS\_" prepended to their names. For example, the `$HOSTNAME$` macro would be available as an environment variable named "NAGIOS\_HOSTNAME".

## **Macro Validity**

Although macros can be used in all commands you define, not all macros may be "valid" in a particular type of command. For example, some macros may only be valid during service notification commands, whereas other may only be valid during host check commands. There are ten types of commands that Nagios recognizes and treats differently. They are as follows:

1. Service checks
2. Service notifications
3. Host checks
4. Host notifications
5. Service [event handlers](#) and/or a global service event handler
6. Host [event handlers](#) and/or a global host event handler
7. [OCS](#)P command
8. [OCH](#)P command
9. Service [performance data](#) commands
10. Host [performance data](#) commands

The tables below list all macros currently available in Nagios, along with a brief description of each and the types of commands in which they are valid. If a macro is used in a command in which it is invalid, it is replaced with an empty string. It should be noted that macros consist of all uppercase characters and are enclosed in \$ characters.

### Macro Availability Chart

#### Legend:

|     |                            |
|-----|----------------------------|
| No  | The macro is not available |
| Yes | The macro is available     |

| Macro Name                             | Service Checks | Service Notifications | Host Checks      | Host Notifications | Service Event Handlers, Global Service Event Handler, <a href="#">OCS</a> P Command | Host Event Handlers, Global Host Event Handler, <a href="#">OCH</a> P Command | Service Performance Data Commands | Host Performance Data Commands |
|----------------------------------------|----------------|-----------------------|------------------|--------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------|-----------------------------------|--------------------------------|
| Host Macros: <sup>3</sup>              |                |                       |                  |                    |                                                                                     |                                                                               |                                   |                                |
| <a href="#">\$HOSTNAMES\$</a>          | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTALIAS\$</a>          | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTADDRESS\$</a>        | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTSTATES\$</a>         | Yes            | Yes                   | Yes <sup>1</sup> | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTSTATEID\$</a>        | Yes            | Yes                   | Yes <sup>1</sup> | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTSTATETYPE\$</a>      | Yes            | Yes                   | Yes <sup>1</sup> | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTATTEMPTS\$</a>       | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTLATENCY\$</a>        | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTEXECUTIONTIMES\$</a> | Yes            | Yes                   | Yes <sup>1</sup> | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTDURATION\$</a>       | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTDURATIONSEC\$</a>    | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTDOWNTIMES\$</a>      | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |
| <a href="#">\$HOSTPERCENTCHANGES\$</a> | Yes            | Yes                   | Yes              | Yes                | Yes                                                                                 | Yes                                                                           | Yes                               | Yes                            |

|                            |                  |                  |                  |                  |     |     |     |     |
|----------------------------|------------------|------------------|------------------|------------------|-----|-----|-----|-----|
| \$HOSTGROUPNAME\$          | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$HOSTGROUPALIAS\$         | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$LASTHOSTCHECK\$          | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$LASTHOSTSTATECHANGE\$    | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$LASTHOSTUP\$             | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$LASTHOSTDOWN\$           | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$LASTHOSTUNREACHABLE\$    | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$HOSTOUTPUT\$             | Yes              | Yes              | Yes <sup>1</sup> | Yes              | Yes | Yes | Yes | Yes |
| \$HOSTPERFDATA\$           | Yes              | Yes              | Yes <sup>1</sup> | Yes              | Yes | Yes | Yes | Yes |
| \$HOSTCHECKCOMMAND\$       | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$HOSTACKAUTHOR\$          | No               | No               | No               | Yes              | No  | No  | No  | No  |
| \$HOSTACKCOMMENT\$         | No               | No               | No               | Yes              | No  | No  | No  | No  |
| \$HOSTACTIONURL\$          | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$HOSTNOTESURL\$           | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| \$HOSTNOTES\$              | Yes              | Yes              | Yes              | Yes              | Yes | Yes | Yes | Yes |
| Service Macros:            |                  |                  |                  |                  |     |     |     |     |
| \$SERVICEDESC\$            | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICESTATE\$           | Yes <sup>2</sup> | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICESTATEID\$         | Yes <sup>2</sup> | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICESTATETYPE\$       | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEATTEMPT\$         | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICELATENCY\$         | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEEXECUTIONTIME\$   | Yes <sup>2</sup> | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEDURATION\$        | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEDURATIONSEC\$     | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEDOWNTIME\$        | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEPERCENTCHANGE\$   | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEGROUPNAME\$       | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEGROUPALIAS\$      | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$LASTSERVICECHECK\$       | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$LASTSERVICESTATECHANGE\$ | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$LASTSERVICEOK\$          | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$LASTSERVICEWARNINGS\$    | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$LASTSERVICEUNKNOWN\$     | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$LASTSERVICECRITICAL\$    | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEOUTPUT\$          | Yes <sup>2</sup> | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEPERFDATA\$        | Yes <sup>2</sup> | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICECHECKCOMMAND\$    | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICEACKAUTHOR\$       | No               | Yes              | No               | No               | No  | No  | No  | No  |
| \$SERVICEACKCOMMENT\$      | No               | Yes              | No               | No               | No  | No  | No  | No  |
| \$SERVICEACTIONURL\$       | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICENOTESURL\$        | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| \$SERVICENOTES\$           | Yes              | Yes              | No               | No               | Yes | No  | Yes | No  |
| Summary Macros:            |                  |                  |                  |                  |     |     |     |     |
| \$TOTALHOSTSUP\$           | Yes              | Yes <sup>4</sup> | Yes              | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |

|                                    |     |                  |     |                  |     |     |     |     |
|------------------------------------|-----|------------------|-----|------------------|-----|-----|-----|-----|
| \$TOTALHOSTSDOWN\$                 | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALHOSTSUNREACHABLE\$          | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALHOSTSDOWNUNHANDLED\$        | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALHOSTSUNREACHABLEUNHANDLED\$ | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALHOSTPROBLEMS\$              | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALHOSTPROBLEMSUNHANDLED\$     | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICESOK\$                | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICESWARNING\$           | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICESCRITICAL\$          | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICESUNKNOWN\$           | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICESWARNINGUNHANDLED\$  | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICESCRITICALUNHANDLED\$ | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICESUNKNOWNUNHANDLED\$  | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICEPROBLEMS\$           | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| \$TOTALSERVICEPROBLEMSUNHANDLED\$  | Yes | Yes <sup>4</sup> | Yes | Yes <sup>4</sup> | Yes | Yes | Yes | Yes |
| Notification Macros:               |     |                  |     |                  |     |     |     |     |
| \$NOTIFICATIONTYPE\$               | No  | Yes              | No  | Yes              | No  | No  | No  | No  |
| \$NOTIFICATIONNUMBERS\$            | No  | Yes              | No  | Yes              | No  | No  | No  | No  |
| Contact Macros:                    |     |                  |     |                  |     |     |     |     |
| \$CONTACTNAME\$                    | No  | Yes              | No  | Yes              | No  | No  | No  | No  |
| \$CONTACTALIAS\$                   | No  | Yes              | No  | Yes              | No  | No  | No  | No  |
| \$CONTACTEMAIL\$                   | No  | Yes              | No  | Yes              | No  | No  | No  | No  |
| \$CONTACTPAGER\$                   | No  | Yes              | No  | Yes              | No  | No  | No  | No  |
| \$CONTACTADDRESSn\$                | No  | Yes              | No  | Yes              | No  | No  | No  | No  |
| Date Macros:                       |     |                  |     |                  |     |     |     |     |
| \$LONGDATETIME\$                   | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$SHORTDATETIME\$                  | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$DATE\$                           | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$TIME\$                           | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$TIMET\$                          | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| File Macros:                       |     |                  |     |                  |     |     |     |     |
| \$MAINCONFIGFILE\$                 | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$STATUSDATAFILE\$                 | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$COMMENTDATAFILE\$                | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$DOWNTIMedataFILE\$               | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$RETENTIONDATAFILE\$              | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$OBJECTCACHEFILE\$                | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$TEMPFILE\$                       | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$LOGFILE\$                        | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$RESOURCEFILE\$                   | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$COMMANDFILE\$                    | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$HOSTPERFDATAFILE\$               | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| \$SERVICEPERFDATAFILE\$            | Yes | Yes              | Yes | Yes              | Yes | Yes | Yes | Yes |
| Misc Macros:                       |     |                  |     |                  |     |     |     |     |

|                      |     |     |     |     |     |     |     |     |
|----------------------|-----|-----|-----|-----|-----|-----|-----|-----|
| \$PROCESSSTARTTIME\$ | Yes | Yes | Yes | Yes | Yes | Yes | Yes | Yes |
| \$ADMINEMAIL\$       | Yes | Yes | Yes | Yes | Yes | Yes | Yes | Yes |
| \$ADMINPAGER\$       | Yes | Yes | Yes | Yes | Yes | Yes | Yes | Yes |
| \$ARGn\$             | Yes | Yes | Yes | Yes | Yes | Yes | Yes | Yes |
| \$USERn\$            | Yes | Yes | Yes | Yes | Yes | Yes | Yes | Yes |

## Macro Descriptions

|                           |                                                                                                                                                                                                                                                                                                                 |  |  |  |  |  |  |  |
|---------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--|--|--|--|--|--|--|
| Host Macros: <sup>3</sup> |                                                                                                                                                                                                                                                                                                                 |  |  |  |  |  |  |  |
| \$HOSTNAME\$              | Short name for the host (i.e. "biglinuxbox"). This value is taken from the <i>host_name</i> directive in the <a href="#">host definition</a> .                                                                                                                                                                  |  |  |  |  |  |  |  |
| \$HOSTALIAS\$             | Long name/description for the host. This value is taken from the <i>alias</i> directive in the <a href="#">host definition</a> .                                                                                                                                                                                |  |  |  |  |  |  |  |
| \$HOSTADDRESS\$           | Address of the host. This value is taken from the <i>address</i> directive in the <a href="#">host definition</a> .                                                                                                                                                                                             |  |  |  |  |  |  |  |
| \$HOSTSTATE\$             | A string indicating the current state of the host ("UP", "DOWN", or "UNREACHABLE").                                                                                                                                                                                                                             |  |  |  |  |  |  |  |
| \$HOSTSTATEID\$           | A number that corresponds to the current state of the host: 0=UP, 1=DOWN, 2=UNREACHABLE.                                                                                                                                                                                                                        |  |  |  |  |  |  |  |
| \$HOSTSTATETYPE\$         | A string indicating the <a href="#">state type</a> for the current host check ("HARD" or "SOFT"). Soft states occur when host checks return a non-OK (non-UP) state and are in the process of being retried. Hard states result when host checks have been checked a specified maximum number of times.         |  |  |  |  |  |  |  |
| \$HOSTATTEMPT\$           | The number of the current host check retry. For instance, if this is the second time that the host is being rechecked, this will be the number two. Current attempt number is really only useful when writing host event handlers for "soft" states that take a specific action based on the host retry number. |  |  |  |  |  |  |  |

|                        |                                                                                                                                                                                                                                                                                                                                           |  |  |  |  |  |  |  |  |
|------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--|--|--|--|--|--|--|--|
| \$HOSTLATENCY\$        | A (floating point) number indicating the number of seconds that a <i>scheduled</i> host check lagged behind its scheduled check time. For instance, if a check was scheduled for 03:14:15 and it didn't get executed until 03:14:17, there would be a check latency of 2.0 seconds. On-demand host checks have a latency of zero seconds. |  |  |  |  |  |  |  |  |
| \$HOSTEXECUTIONTIME\$  | A (floating point) number indicating the number of seconds that the host check took to execute (i.e. the amount of time the check was executing).                                                                                                                                                                                         |  |  |  |  |  |  |  |  |
| \$HOSTDURATION\$       | A string indicating the amount of time that the host has spent in its current state. Format is "XXh YYm ZZs", indicating hours, minutes and seconds.                                                                                                                                                                                      |  |  |  |  |  |  |  |  |
| \$HOSTDURATIONSEC\$    | A number indicating the number of seconds that the host has spent in its current state.                                                                                                                                                                                                                                                   |  |  |  |  |  |  |  |  |
| \$HOSTDOWNTIME\$       | A number indicating the current "downtime depth" for the host. If this host is currently in a period of <a href="#">scheduled downtime</a> , the value will be greater than zero. If the host is not currently in a period of downtime, this value will be zero.                                                                          |  |  |  |  |  |  |  |  |
| \$HOSTPERCENTCHANGES\$ | A (floating point) number indicating the percent state change the host has undergone. Percent state change is used by the <a href="#">flap detection</a> algorithm.                                                                                                                                                                       |  |  |  |  |  |  |  |  |
| \$HOSTGROUPNAME\$      | The short name of the hostgroup that this host belongs to. This value is taken from the <i>hostgroup_name</i> directive in the <a href="#">hostgroup definition</a> . If the host belongs to more than one hostgroup this macro will contain the name of just one of them.                                                                |  |  |  |  |  |  |  |  |

|                         |                                                                                                                                                                                                                                                                        |  |  |  |  |  |  |  |  |
|-------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--|--|--|--|--|--|--|--|
| \$HOSTGROUPALIAS\$      | The longer name/alias of the hostgroup that this host belongs to. This value is taken from the <i>alias</i> directive in the <a href="#">hostgroup definition</a> . If the host belongs to more than one hostgroup, this macro contains the alias of just one of them. |  |  |  |  |  |  |  |  |
| \$LASTHOSTCHECK\$       | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which a check of the host was last performed.                                                                                                                               |  |  |  |  |  |  |  |  |
| \$LASTHOSTSTATECHANGE\$ | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time the host last changed state.                                                                                                                                                   |  |  |  |  |  |  |  |  |
| \$LASTHOSTUP\$          | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which the host was last detected as being in an UP state.                                                                                                                   |  |  |  |  |  |  |  |  |
| \$LASTHOSTDOWN\$        | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which the host was last detected as being in a DOWN state.                                                                                                                  |  |  |  |  |  |  |  |  |
| \$LASTHOSTUNREACHABLE\$ | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which the host was last detected as being in an UNREACHABLE state.                                                                                                          |  |  |  |  |  |  |  |  |
| \$HOSTOUTPUT\$          | The text output from the last host check (i.e. "Ping OK").                                                                                                                                                                                                             |  |  |  |  |  |  |  |  |
| \$HOSTPERFDATA\$        | This macro contains any <a href="#">performance data</a> that may have been returned by the last host check.                                                                                                                                                           |  |  |  |  |  |  |  |  |
| \$HOSTCHECKCOMMAND\$    | This macro contains the name of the command (along with any arguments passed to it) used to perform the host check.                                                                                                                                                    |  |  |  |  |  |  |  |  |





|                          |                                                                                                                                                                                                                                                                                                                             |  |  |  |  |  |  |  |  |
|--------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--|--|--|--|--|--|--|--|
| \$SERVICESTATETYPE\$     | A string indicating the <a href="#">state type</a> for the current service check ("HARD" or "SOFT"). Soft states occur when service checks return a non-OK state and are in the process of being retried. Hard states result when service checks have been checked a specified maximum number of times.                     |  |  |  |  |  |  |  |  |
| \$SERVICEATTEMPT\$       | The number of the current service check retry. For instance, if this is the second time that the service is being rechecked, this will be the number two. Current attempt number is really only useful when writing service event handlers for "soft" states that take a specific action based on the service retry number. |  |  |  |  |  |  |  |  |
| \$SERVICELATENCY\$       | A (floating point) number indicating the number of seconds that a scheduled service check lagged behind its scheduled check time. For instance, if a check was scheduled for 03:14:15 and it didn't get executed until 03:14:17, there would be a check latency of 2.0 seconds.                                             |  |  |  |  |  |  |  |  |
| \$SERVICEEXECUTIONTIME\$ | A (floating point) number indicating the number of seconds that the service check took to execute (i.e. the amount of time the check was executing).                                                                                                                                                                        |  |  |  |  |  |  |  |  |
| \$SERVICEDURATION\$      | A string indicating the amount of time that the service has spent in its current state. Format is "XXh YYm ZZs", indicating hours, minutes and seconds.                                                                                                                                                                     |  |  |  |  |  |  |  |  |
| \$SERVICEDURATIONSEC\$   | A number indicating the number of seconds that the service has spent in its current state.                                                                                                                                                                                                                                  |  |  |  |  |  |  |  |  |

|                            |                                                                                                                                                                                                                                                                                             |  |  |  |  |  |  |  |  |
|----------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--|--|--|--|--|--|--|--|
| \$SERVICEDOWNTIME\$        | A number indicating the current "downtime depth" for the service. If this service is currently in a period of <a href="#">scheduled downtime</a> , the value will be greater than zero. If the service is not currently in a period of downtime, this value will be zero.                   |  |  |  |  |  |  |  |  |
| \$SERVICEPERCENTCHANGE\$   | A (floating point) number indicating the percent state change the service has undergone. Percent state change is used by the <a href="#">flap detection</a> algorithm.                                                                                                                      |  |  |  |  |  |  |  |  |
| \$SERVICEGROUPNAME\$       | The short name of the servicegroup that this service belongs to. This value is taken from the <i>servicegroup_name</i> directive in the <a href="#">servicegroup</a> definition. If the service belongs to more than one servicegroup this macro will contain the name of just one of them. |  |  |  |  |  |  |  |  |
| \$SERVICEGROUPALIAS\$      | The long name/alias of the servicegroup that this service belongs to. This value is taken from the <i>alias</i> directive in the <a href="#">servicegroup</a> definition. If the service belongs to more than one servicegroup this macro will contain the name of just one of them.        |  |  |  |  |  |  |  |  |
| \$LASTSERVICECHECK\$       | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which a check of the service was last performed.                                                                                                                                                 |  |  |  |  |  |  |  |  |
| \$LASTSERVICESTATECHANGE\$ | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time the service last changed state.                                                                                                                                                                     |  |  |  |  |  |  |  |  |
| \$LASTSERVICEOK\$          | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which the service was last detected as being in an OK state.                                                                                                                                     |  |  |  |  |  |  |  |  |

|                         |                                                                                                                                                                                                                                |  |  |  |  |  |  |  |  |
|-------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--|--|--|--|--|--|--|--|
| \$LASTSERVICEWARNING\$  | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which the service was last detected as being in a WARNING state.                                                                    |  |  |  |  |  |  |  |  |
| \$LASTSERVICEUNKNOWN\$  | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which the service was last detected as being in an UNKNOWN state.                                                                   |  |  |  |  |  |  |  |  |
| \$LASTSERVICECRITICAL\$ | This is a timestamp in time_t format (seconds since the UNIX epoch) indicating the time at which the service was last detected as being in a CRITICAL state.                                                                   |  |  |  |  |  |  |  |  |
| \$SERVICEOUTPUT\$       | The text output from the last service check (i.e. "Ping OK").                                                                                                                                                                  |  |  |  |  |  |  |  |  |
| \$SERVICEPERFDATA\$     | This macro contains any <a href="#">performance data</a> that may have been returned by the last service check.                                                                                                                |  |  |  |  |  |  |  |  |
| \$SERVICECHECKCOMMAND\$ | This macro contains the name of the command (along with any arguments passed to it) used to perform the service check.                                                                                                         |  |  |  |  |  |  |  |  |
| \$SERVICEACKAUTHOR\$    | A string containing the name of the user who acknowledged the service problem. This macro is only valid in notifications where the \$NOTIFICATIONTYPE\$ macro is set to "ACKNOWLEDGEMENT".                                     |  |  |  |  |  |  |  |  |
| \$SERVICEACKCOMMENT\$   | A string containing the acknowledgement comment that was entered by the user who acknowledged the service problem. This macro is only valid in notifications where the \$NOTIFICATIONTYPE\$ macro is set to "ACKNOWLEDGEMENT". |  |  |  |  |  |  |  |  |
| \$SERVICEACTIONURL\$    | Action URL for the service. This value is taken from the <i>action_url</i> directive in the <a href="#">extended service information definition</a> .                                                                          |  |  |  |  |  |  |  |  |

|                             |                                                                                                                                                                                                                                                                                                                                                                                                        |  |  |  |  |  |  |  |  |
|-----------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--|--|--|--|--|--|--|--|
| \$SERVICENOTESURL\$         | Notes URL for the service. This value is taken from the <i>notes_url</i> directive in the <a href="#">extended service information definition</a> .                                                                                                                                                                                                                                                    |  |  |  |  |  |  |  |  |
| \$SERVICENOTES\$            | Notes for the service. This value is taken from the <i>notes</i> directive in the <a href="#">extended service information definition</a> .                                                                                                                                                                                                                                                            |  |  |  |  |  |  |  |  |
| Notification Macros:        |                                                                                                                                                                                                                                                                                                                                                                                                        |  |  |  |  |  |  |  |  |
| \$NOTIFICATIONTYPE\$        | A string identifying the type of notification that is being sent ("PROBLEM", "RECOVERY", "ACKNOWLEDGEMENT", "FLAPPINGSTART" or "FLAPPINGSTOP").                                                                                                                                                                                                                                                        |  |  |  |  |  |  |  |  |
| \$NOTIFICATIONNUMBER\$      | The current notification number for the service or host. The notification number increases by one (1) each time a new notification is sent out for a host or service (except for acknowledgements). The notification number is reset to 0 when the host or service recovers ( <i>after</i> the recovery notification has gone out). Acknowledgements do not cause the notification number to increase. |  |  |  |  |  |  |  |  |
| SUMMARY Macros:             |                                                                                                                                                                                                                                                                                                                                                                                                        |  |  |  |  |  |  |  |  |
| \$TOTALHOSTSUP\$            | This macro reflects the total number of hosts that are currently in an UP state.                                                                                                                                                                                                                                                                                                                       |  |  |  |  |  |  |  |  |
| \$TOTALHOSTSDOWN\$          | This macro reflects the total number of hosts that are currently in a DOWN state.                                                                                                                                                                                                                                                                                                                      |  |  |  |  |  |  |  |  |
| \$TOTALHOSTSUNREACHABLE\$   | This macro reflects the total number of hosts that are currently in an UNREACHABLE state.                                                                                                                                                                                                                                                                                                              |  |  |  |  |  |  |  |  |
| \$TOTALHOSTSDOWNUNHANDLED\$ | This macro reflects the total number of hosts that are currently in a DOWN state that are not currently being "handled". Unhandled host problems are those that are not acknowledged, are not currently in scheduled downtime, and for which checks are currently enabled.                                                                                                                             |  |  |  |  |  |  |  |  |









|                         |                                                                                                                                                                                                                                                                                             |  |  |  |  |  |  |  |  |
|-------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--|--|--|--|--|--|--|--|
| \$OBJECTCACHEFILE\$     | The location of the <a href="#">object cache file</a> .                                                                                                                                                                                                                                     |  |  |  |  |  |  |  |  |
| \$TEMPFILE\$            | The location of the <a href="#">temp file</a> .                                                                                                                                                                                                                                             |  |  |  |  |  |  |  |  |
| \$LOGFILE\$             | The location of the <a href="#">log file</a> .                                                                                                                                                                                                                                              |  |  |  |  |  |  |  |  |
| \$RESOURCEFILE\$        | The location of the <a href="#">resource file</a> .                                                                                                                                                                                                                                         |  |  |  |  |  |  |  |  |
| \$COMMANDFILE\$         | The location of the <a href="#">command file</a> .                                                                                                                                                                                                                                          |  |  |  |  |  |  |  |  |
| \$HOSTPERFDATAFILE\$    | The location of the host performance data file (if defined).                                                                                                                                                                                                                                |  |  |  |  |  |  |  |  |
| \$SERVICEPERFDATAFILE\$ | The location of the service performance data file (if defined).                                                                                                                                                                                                                             |  |  |  |  |  |  |  |  |
| Misc Macros:            |                                                                                                                                                                                                                                                                                             |  |  |  |  |  |  |  |  |
| \$PROCESSSTARTTIME\$    | Time stamp in time_t format (seconds since the UNIX epoch) indicating when the Nagios process was last (re)started. You can determine the number of seconds that Nagios has been running (since it was last restarted) by subtracting \$PROCESSSTARTTIME\$ from <a href="#">\$TIMET\$</a> . |  |  |  |  |  |  |  |  |
| \$ADMINEMAIL\$          | Global administrative email address. This value is taken from the <a href="#">admin_email</a> directive.                                                                                                                                                                                    |  |  |  |  |  |  |  |  |
| \$ADMINPAGER\$          | Global administrative pager number/address. This value is taken from the <a href="#">admin_pager</a> directive.                                                                                                                                                                             |  |  |  |  |  |  |  |  |
| \$ARGn\$                | The <i>n</i> th argument passed to the command (notification, event handler, service check, etc.). Nagios supports up to 32 argument macros (\$ARG1\$ through \$ARG32\$).                                                                                                                   |  |  |  |  |  |  |  |  |
| \$USERn\$               | The <i>n</i> th user-definable macro. User macros can be defined in one or more <a href="#">resource files</a> . Nagios supports up to 32 user macros (\$USER1\$ through \$USER32\$).                                                                                                       |  |  |  |  |  |  |  |  |

## Notes

<sup>1</sup> These macros are not valid for the host they are associated with when that host is being checked (i.e. they make no sense, as they haven't been determined yet).

<sup>2</sup> These macros are not valid for the service they are associated with when that service is being checked (i.e. they make no sense, as they haven't been determined yet).

<sup>3</sup> When host macros are used in service-related commands (i.e. service notifications, event handlers, etc) they refer to the host that the service is associated with.

<sup>4</sup> When host and service summary macros are used in notification commands, the totals are filtered to reflect only those hosts and services for which the contact is authorized (i.e. hosts and services they are configured to receive notifications for).

---

# Information On The CGIs

---

## Introduction

The various CGIs distributed with Nagios are described here, along with the authorization requirements for accessing and using each CGI. By default the CGIs require that you have authenticated to the web server and are authorized to view any information you are requesting. For more information on configuring your web server and CGI configuration file to allow for this, read the sections on [setting up the web interface](#) and [CGI authorization](#).

## Index

- Status CGI
- Status map CGI
- WAP interface CGI
- Status world CGI (VRML)
- Tactical overview CGI
- Network outages CGI
- Configuration CGI
- Command CGI
- Extended information CGI
- Event log CGI
- Alert history CGI
- Notifications CGI
- Trends CGI
- Availability reporting CGI
- Alert histogram CGI
- Alert summary CGI

## Status CGI



File Name: **status.cgi**

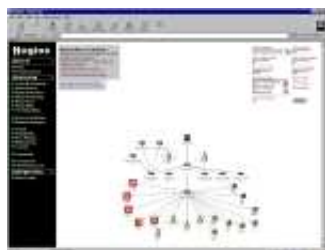
**Description:**

This is the most important CGI included with Nagios. It allows you to view the current status of all hosts and services that are being monitored. The status CGI can produce two main types of output - a status overview of all host groups (or a particular host group) and a detailed view of all services (or those associated with a particular host). Pretty icons can be associated with hosts by defining [extended host and service information](#) entries.

**Authorization Requirements:**

- If you are [authorized for all hosts](#) you can view all hosts **and** all services.
- If you are [authorized for all services](#) you can view all services.
- If you are an *authenticated contact* you can view all hosts and services for which you are a contact.

## Status Map CGI



File Name: **statusmap.cgi**

**Description:**

This CGI creates a map of all hosts that you have defined on your network. The CGI uses Thomas Boutell's [gd](#) library (version 1.6.3 or higher) to create a PNG image of your network layout. The coordinates used when drawing each host (along with the optional pretty icons) are taken from [extended host information](#) definitions. If you'd prefer to let the CGI automatically generate drawing coordinates for you, use the [default\\_statusmap\\_layout](#) directive to specify a layout algorithm that should be used.

**Authorization Requirements:**

- If you are [authorized for all hosts](#) you can view all hosts.
- If you are an *authenticated contact* you can view hosts for which you are a contact.

Note: Users who are not authorized to view specific hosts will see *unknown* nodes in those positions. I realize that they really shouldn't see *anything* there, but it doesn't make sense to even generate the map if you can't see all the host dependencies...

## WAP Interface CGI



File Name: **statuswml.cgi**

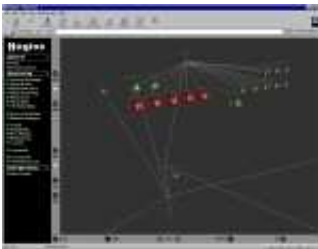
### Description:

This CGI serves as a WAP interface to network status information. If you have a WAP-enabled device (i.e. an Internet-ready cellphone), you can view status information while you're on the go. Different status views include hostgroup summary, hostgroup overview, host detail, service detail, all problems, and unhandled problems. In addition to viewing status information, you can also disable notifications and checks and acknowledge problems from your cellphone. Pretty cool, huh?

### Authorization Requirements:

- If you are *authorized for system information* you can view Nagios process information.
- If you are *authorized for all hosts* you can view status data for all hosts **and** services.
- If you are *authorized for all services* you can view status data for all services.
- If you are an *authenticated contact* you can view status data for all hosts and services for which you are a contact.

### Status World CGI (VRML)



File Name: **statuswrl.cgi**

**Description:**

This CGI creates a 3-D VRML model of all hosts that you have defined on your network. Coordinates used when drawing the hosts (as well as pretty texture maps) are defined using [extended host information](#) definitions. If you'd prefer to let the CGI automatically generate drawing coordinates for you, use the [default\\_statuswrl\\_layout](#) directive to specify a layout algorithm that should be used. You'll need a VRML browser (like [Cortona](#), [Cosmo Player](#) or [WorldView](#)) installed on your system before you can actually view the generated model.

**Authorization Requirements:**

- If you are [authorized for all hosts](#) you can view all hosts.
- If you are an *authenticated contact* you can view hosts for which you are a contact.

Note: Users who are not authorized to view specific hosts will see *unknown* nodes in those positions. I realize that they really shouldn't see *anything* there, but it doesn't make sense to even generate the map if you can't see all the host dependencies...

**Tactical Overview CGI**

File Name: **tac.cgi**

**Description:**

This CGI is designed to server as a "birds-eye view" of all network monitoring activity. It allows you to quickly see network outages, host status, and service status. It distinguishes between problems that have been "handled" in some way (i.e. been acknowledged, had notifications disabled, etc.) and those which have not been handled, and thus need attention. Very useful if you've got a lot of hosts/services you're monitoring and you need to keep a single screen up to alert you of problems.

**Authorization Requirements:**

- If you are [authorized for all hosts](#) you can view all hosts **and** all services.
- If you are [authorized for all services](#) you can view all services.
- If you are an *authenticated contact* you can view all hosts and services for which you are a contact.

**Network Outages CGI**



File Name: **outages.cgi**

**Description:**

This CGI will produce a listing of "problem" hosts on your network that are causing network outages. This can be particularly useful if you have a large network and want to quickly identify the source of the problem. Hosts are sorted based on the severity of the outage they are causing. More information on how the network outage CGI works can be found [here](#).

**Authorization Requirements:**

- If you are *authorized for all hosts* you can view all hosts.
- If you are an *authenticated contact* you can view hosts for which you are a contact.

## Configuration CGI



File Name: **config.cgi**

**Description:**

This CGI allows you to view objects (i.e. hosts, host groups, contacts, contact groups, time periods, services, etc.) that you have defined in your [object configuration file\(s\)](#).

**Authorization Requirements:**

- You must be *authorized for configuration information* in order to any kind of configuration information.

## Command CGI



File Name: **cmd.cgi**

### Description:

This CGI allows you to send commands to the Nagios process. Although this CGI has several arguments, you would be better to leave them alone. Most will change between different revisions of Nagios. Use the [extended information CGI](#) as a starting point for issuing commands.

### Authorization Requirements:

- You must be *authorized for system commands* in order to issue commands that affect the Nagios process (restarts, shutdowns, mode changes, etc.).
- If you are *authorized for all host commands* you can issue commands for all hosts **and** services.
- If you are *authorized for all service commands* you can issue commands for all services.
- If you are an *authenticated contact* you can issue commands for all hosts and services for which you are a contact.

### Notes:

- If you have chosen not to [use authentication](#) with the CGIs, this CGI will *not* allow anyone to issue commands to Nagios. This is done for your own protection. I would suggest removing this CGI altogether if you decide not to use authentication with the CGIs.

## Extended Information CGI



File Name: **extinfo.cgi**



**Description:**

This CGI allows you to view Nagios process information, host and service state statistics, host and service comments, and more. It also serves as a launching point for sending commands to Nagios via the [command CGI](#). Although this CGI has several arguments, you would be better to leave them alone - they are likely to change between different releases of Nagios. You can access this CGI by clicking on the 'Network Health' and 'Process Information' links on the side navigation bar, or by clicking on a host or service link in the output of the [status CGI](#).

**Authorization Requirements:**

- You must be [authorized for system information](#) in order to view Nagios process information.
- If you are [authorized for all hosts](#) you can view extended information for all hosts **and** services.
- If you are [authorized for all services](#) you can view extended information for all services.
- If you are an *authenticated contact* you can view extended information for all hosts and services for which you are a contact.

## Event Log CGI



File Name: **showlog.cgi**

**Description:**

This CGI will display the [log file](#). If you have [log rotation](#) enabled, you can browse notifications present in archived log files by using the navigational links near the top of the page.

**Authorization Requirements:**

- You must be [authorized for system information](#) in order to view the log file.

## Alert History CGI



File Name: **history.cgi**

### Description:

This CGI is used to display the history of problems with either a particular host or all hosts. The output is basically a subset of the information that is displayed by the [log file CGI](#). You have the ability to filter the output to display only the specific types of problems you wish to see (i.e. hard and/or soft alerts, various types of service and host alerts, all types of alerts, etc.). If you have [log rotation](#) enabled, you can browse history information present in archived log files by using the navigational links near the top of the page.

### Authorization Requirements:

- If you are [authorized for all hosts](#) you can view history information for all hosts **and** all services.
- If you are [authorized for all services](#) you can view history information for all services.
- If you are an *authenticated contact* you can view history information for all services and hosts for which you are a contact.

## Notifications CGI



File Name: **notifications.cgi**

### Description:

This CGI is used to display host and service notifications that have been sent to various contacts. The output is basically a subset of the information that is displayed by the [log file CGI](#). You have the ability to filter the output to display only the specific types of notifications you wish to see (i.e. service notifications, host notifications, notifications sent to specific contacts, etc). If you have [log rotation](#) enabled, you can browse notifications present in archived log files by using the navigational links near the top of the page.

### Authorization Requirements:

- If you are [authorized for all hosts](#) you can view notifications for all hosts **and** all services.
- If you are [authorized for all services](#) you can view notifications for all services.
- If you are an *authenticated contact* you can view notifications for all services and hosts for which you are a contact.

## Trends CGI



File Name: **trends.cgi**

#### Description:

This CGI is used to create a graph of host or service states over an arbitrary period of time. In order for this CGI to be of much use, you should enable [log rotation](#) and keep archived logs in the path specified by the [log\\_archive\\_path](#) directive. The CGI uses Thomas Boutell's [gd](#) library (version 1.6.3 or higher) to create the trends image.

#### Authorization Requirements:

- If you are [authorized for all hosts](#) you can view trends for all hosts **and** all services.
- If you are [authorized for all services](#) you can view trends for all services.
- If you are an *authenticated contact* you can view trends for all services and hosts for which you are a contact.

### Availability Reporting CGI



File Name: **avail.cgi**

#### Description:

This CGI is used to report on the availability of hosts and services over a user-specified period of time. In order for this CGI to be of much use, you should enable [log rotation](#) and keep archived logs in the path specified by the [log\\_archive\\_path](#) directive.

#### Authorization Requirements:

- If you are [authorized for all hosts](#) you can view availability data for all hosts **and** all services.
- If you are [authorized for all services](#) you can view availability data for all services.
- If you are an *authenticated contact* you can view availability data for all services and hosts for which you are a contact.

## Alert Histogram CGI



File Name: **histogram.cgi**

### Description:

This CGI is used to report on the availability of hosts and services over a user-specified period of time. In order for this CGI to be of much use, you should enable [log rotation](#) and keep archived logs in the path specified by the [log\\_archive\\_path](#) directive. The CGI uses Thomas Boutell's [gd](#) library (version 1.6.3 or higher) to create the histogram image.

### Authorization Requirements:

- If you are [authorized for all hosts](#) you can view histograms for all hosts **and** all services.
- If you are [authorized for all services](#) you can view histograms for all services.
- If you are an [authenticated contact](#) you can view histograms for all services and hosts for which you are a contact.

## Alert Summary CGI



File Name: **summary.cgi**

### Description:

This CGI provides some generic reports about host and service alert data, including alert totals, top alert producers, etc.

### Authorization Requirements:

- If you are [authorized for all hosts](#) you can view summary information for all hosts **and** all services.
- If you are [authorized for all services](#) you can view summary information for all services.
- If you are an [authenticated contact](#) you can view summary information for all services and hosts for which you are a contact.



# Custom CGI Headers and Footers

---

## Introduction

If you're doing custom installs of Nagios for clients, you may want to have a custom header and/or footer displayed in the output of the [CGIs](#). This is particularly useful for displaying support contact information, etc. to the end user.

It is important to note that, unless the custom header and footer files are executable, they are **not** processed in any way before they are displayed. The contents of the header and footer include files are simply read and displayed in the CGI output. That means they can only contain information a web browser can understand (HTML, JavaScript, etc.).

If the custom header and footer files are executable, then the files are executed and their output returned to the user, so they should output valid HTML. Using this you can run your own custom designed CGI to insert data into the nagios display. This has been used to insert graphs from rrdtool using ddraw and command menus into the nagios display pane. The executable customer header and footer files are run with the same CGI environment as the main nagios cgi, so your files can parse the query information, authenticated user information etc. to produce appropriate output.

## How Does It Work?

You can include custom headers and footers in the output of the CGIs by dropping some appropriately named HTML files in the *ssi/* subdirectory of the Nagios HTML directory (i.e. */usr/local/nagios/share/ssi*).

Custom headers are included immediately after the **<BODY>** tag in the CGI output, while custom footers are included immediately before the closing **</BODY>** tag.

There are two types of customer headers and footers:

- **Global headers/footers.** These files should be named *common-header.ssi* and *common-footer.ssi*, respectively. If these files exist, they will be included in the output of all CGIs.
- **CGI-specific headers/footers.** These files should be named in the format *CGINAME-header.ssi* and *CGINAME-footer.ssi*, where *CGINAME* is the physical name of the CGI without the *.cgi* extension. For example, the header and footer files for the [alert summary CGI](#) (*summary.cgi*) would be named *summary-header.ssi* and *summary-footer.ssi*, respectively.

You are not required to use any custom headers or footers. You can use only a global header if you wish. You can use only CGI-specific headers and a global footer if you wish. Whatever you want. Really.

---

# Template-Based Object Configuration

---

## Notes

When creating and/or editing configuration files, keep the following in mind:

1. Lines that start with a '#' character are taken to be comments and are not processed
2. Directive names are case-sensitive

## Introduction

One of the benefits of using the template-based config file format is that you can create object definitions that have some of their properties inherited from other object definitions. The notion of object inheritance, along with documentation on how to do it, is described [here](#). I strongly suggest that you familiarize yourself with object inheritance once you read over the documentation presented below, as inheritance will make the job of creating and maintaining object definitions much easier than it otherwise would be.

## Time-Saving Tricks

There are several things you can do with template-based object definitions that allow you to create large numbers of objects using just a small number of definitions in your config file(s). One example of such a trick is the ability to define a single service object that creates a service for multiple hosts and/or hostgroups. These tricks are described [here](#).

## Retention Notes

It is important to point out that several directives in host and service definitions may not be picked up by Nagios when you change them. Host and service directives that can exhibit this behavior are marked with an asterisk (\*). The reason for this behavior is due to the fact that Nagios chooses to honor values stored in the [state retention file](#) over values found in the config files, assuming you have [state retention](#) enabled on a program-wide basis *and* the value of the directive is changed during runtime (by submitting an [external command](#)).

One way to get around this problem is to disable the retention of non-status information using the *retain\_nonstatus\_information* directive in the host and service definitions. Disabling this directive will cause Nagios to take the initial values for these directives from your config files, rather than from the state retention file when it (re)starts. Using this option is not recommended, as it may result in some unexpected (from your point of view) results.

Alternatively, you can issue the appropriate [external command](#) or change the value of the host or service directive via the web interface, so that it matches what you've changed it to in the config files. This is usually done by using the [extended information CGI](#). This option takes a bit more work, but is preferable to disabling the retention of non-status information (mentioned above).

## Sample Configuration

A few sample object configuration files are created when you run the configure script - you can find them in the *sample-config/template-object/* subdirectory of the Nagios distribution.

## Object Types

Host definitions  
Host group definitions  
Service definitions  
Service group definitions  
Contact definitions  
Contact group definitions  
Time period definitions  
Command definitions  
Service dependency definitions  
Service escalation definitions  
Host dependency definitions  
Host escalation definitions  
Extended host information definitions  
Extended service information definitions

## Host Definition

### Description:

A host definition is used to define a physical server, workstation, device, etc. that resides on your network.

### Definition Format:

Note: Directives in red are required, while those in black are optional.



```

define host{
    host_name                host_name
    alias                    alias
    address                  address
    parents                  host_names
    hostgroups               hostgroup_names
    check_command            command_name
    max_check_attempts       #
    check_interval           #
    active_checks_enabled    [0/1]
    passive_checks_enabled   [0/1]
    check_period             timeperiod_name
    obsess_over_host         [0/1]
    check_freshness          [0/1]
    freshness_threshold      #
    event_handler            command_name
    event_handler_enabled    [0/1]
    low_flap_threshold       #
    high_flap_threshold      #
    flap_detection_enabled    [0/1]
    process_perf_data        [0/1]
    retain_status_information [0/1]
    retain_nonstatus_information [0/1]
    contact_groups           contact_groups
    notification_interval     #
    notification_period       timeperiod_name
    notification_options      [d,u,r,f]
    notifications_enabled     [0/1]
    stalking_options          [o,d,u]
}

```

Example Definition:

```

define host{
    host_name          bogus-router
    alias              Bogus Router #1
    address            192.168.1.254
    parents            server-backbone
    check_command      check-host-alive
    max_check_attempts 5
    check_period       24x7
    process_perf_data  0
    retain_nonstatus_information 0
    contact_groups     router-admins
    notification_interval 30
    notification_period 24x7
    notification_options d,u,r
}

```

## Directive Descriptions:

- host\_name:** This directive is used to define a short name used to identify the host. It is used in host group and service definitions to reference this particular host. Hosts can have multiple services (which are monitored) associated with them. When used properly, the \$HOSTNAME\$ [macro](#) will contain this short name.
- alias:** This directive is used to define a longer name or description used to identify the host. It is provided in order to allow you to more easily identify a particular host. When used properly, the \$HOSTALIAS\$ [macro](#) will contain this alias/description.
- address:** This directive is used to define the address of the host. Normally, this is an IP address, although it could really be anything you want (so long as it can be used to check the status of the host). You can use a FQDN to identify the host instead of an IP address, but if DNS services are not available this could cause problems. When used properly, the \$HOSTADDRESS\$ [macro](#) will contain this address. **Note:** If you do not specify an address directive in a host definition, the name of the host will be used as its address. A word of caution about doing this, however - if DNS fails, most of your service checks will fail because the plugins will be unable to resolve the host name.
- parents:** This directive is used to define a comma-delimited list of short names of the "parent" hosts for this particular host. Parent hosts are typically routers, switches, firewalls, etc. that lie between the monitoring host and a remote hosts. A router, switch, etc. which is closest to the remote host is considered to be that host's "parent". Read the "Determining Status and Reachability of Network Hosts" document located [here](#) for more information. If this host is on the same network segment as the host doing the monitoring (without any intermediate routers, etc.) the host is considered to be on the local network and will not have a parent host. Leave this value blank if the host does not have a parent host (i.e. it is on the same segment as the Nagios host). The order in which you specify parent hosts has no effect on how things are monitored.

|                                  |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
|----------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>hostgroups:</b>               | This directive is used to identify the <i>short name(s)</i> of the <a href="#">hostgroup(s)</a> that the host belongs to. Multiple hostgroups should be separated by commas. This directive may be used as an alternative to (or in addition to) using the <i>members</i> directive in <a href="#">hostgroup</a> definitions.                                                                                                                                                                                                                                                                                                                                                                                                |
| <b>check_command:</b>            | This directive is used to specify the <i>short name</i> of the <a href="#">command</a> that should be used to check if the host is up or down. Typically, this command would try and ping the host to see if it is "alive". The command must return a status of OK (0) or Nagios will assume the host is down. If you leave this argument blank, the host will <i>not</i> be checked - Nagios will always assume the host is up. This is useful if you are monitoring printers or other devices that are frequently turned off. The maximum amount of time that the notification command can run is controlled by the <a href="#">host_check_timeout</a> option.                                                             |
| <b>max_check_attempts:</b>       | This directive is used to define the number of times that Nagios will retry the host check command if it returns any state other than an OK state. Setting this value to 1 will cause Nagios to generate an alert without retrying the host check again. Note: If you do not want to check the status of the host, you must still set this to a minimum value of 1. To bypass the host check, just leave the <i>check_command</i> option blank.                                                                                                                                                                                                                                                                              |
| <b>check_interval:</b>           | <i>NOTE: Do NOT enable regularly scheduled checks of a host unless you absolutely need to! Host checks are already performed on-demand when necessary, so there are few times when regularly scheduled checks would be needed. Regularly scheduled host checks can negatively impact performance - see the <a href="#">performance tuning tips</a> for more information.</i> This directive is used to define the number of "time units" between regularly scheduled checks of the host. Unless you've changed the <a href="#">interval_length</a> directive from the default value of 60, this number will mean minutes. More information on this value can be found in the <a href="#">check scheduling</a> documentation. |
| <b>active_checks_enabled *:</b>  | This directive is used to determine whether or not active checks (either regularly scheduled or on-demand) of this host are enabled. Values: 0 = disable active host checks, 1 = enable active host checks.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |
| <b>passive_checks_enabled *:</b> | This directive is used to determine whether or not passive checks are enabled for this host. Values: 0 = disable passive host checks, 1 = enable passive host checks.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
| <b>check_period:</b>             | This directive is used to specify the short name of the <a href="#">time period</a> during which active checks of this host can be made.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
| <b>obsess_over_host *:</b>       | This directive determines whether or not checks for the host will be "obsessed" over using the <a href="#">ochp_command</a> .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| <b>check_freshness *:</b>        | This directive is used to determine whether or not <a href="#">freshness checks</a> are enabled for this host. Values: 0 = disable freshness checks, 1 = enable freshness checks.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |

|                                      |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       |
|--------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>freshness_threshold:</b>          | This directive is used to specify the freshness threshold (in seconds) for this host. If you set this directive to a value of 0, Nagios will determine a freshness threshold to use automatically.                                                                                                                                                                                                                                                                                    |
| <b>event_handler:</b>                | This directive is used to specify the <i>short name</i> of the <a href="#">command</a> that should be run whenever a change in the state of the host is detected (i.e. whenever it goes down or recovers). Read the documentation on <a href="#">event handlers</a> for a more detailed explanation of how to write scripts for handling events. The maximum amount of time that the event handler command can run is controlled by the <a href="#">event_handler_timeout</a> option. |
| <b>event_handler_enabled *:</b>      | This directive is used to determine whether or not the event handler for this host is enabled. Values: 0 = disable host event handler, 1 = enable host event handler.                                                                                                                                                                                                                                                                                                                 |
| <b>low_flap_threshold:</b>           | This directive is used to specify the low state change threshold used in flap detection for this host. More information on flap detection can be found <a href="#">here</a> . If you set this directive to a value of 0, the program-wide value specified by the <a href="#">low_host_flap_threshold</a> directive will be used.                                                                                                                                                      |
| <b>high_flap_threshold:</b>          | This directive is used to specify the high state change threshold used in flap detection for this host. More information on flap detection can be found <a href="#">here</a> . If you set this directive to a value of 0, the program-wide value specified by the <a href="#">high_host_flap_threshold</a> directive will be used.                                                                                                                                                    |
| <b>flap_detection_enabled *:</b>     | This directive is used to determine whether or not flap detection is enabled for this host. More information on flap detection can be found <a href="#">here</a> . Values: 0 = disable host flap detection, 1 = enable host flap detection.                                                                                                                                                                                                                                           |
| <b>process_perf_data *:</b>          | This directive is used to determine whether or not the processing of performance data is enabled for this host. Values: 0 = disable performance data processing, 1 = enable performance data processing.                                                                                                                                                                                                                                                                              |
| <b>retain_status_information:</b>    | This directive is used to determine whether or not status-related information about the host is retained across program restarts. This is only useful if you have enabled state retention using the <a href="#">retain_state_information</a> directive. Value: 0 = disable status information retention, 1 = enable status information retention.                                                                                                                                     |
| <b>retain_nonstatus_information:</b> | This directive is used to determine whether or not non-status information about the host is retained across program restarts. This is only useful if you have enabled state retention using the <a href="#">retain_state_information</a> directive. Value: 0 = disable non-status information retention, 1 = enable non-status information retention.                                                                                                                                 |
| <b>contact_groups:</b>               | This is a list of the <i>short names</i> of the <a href="#">contact groups</a> that should be notified whenever there are problems (or recoveries) with this host. Multiple contact groups should be separated by commas.                                                                                                                                                                                                                                                             |

|                                 |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|---------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>notification_interval:</b>   | This directive is used to define the number of "time units" to wait before re-notifying a contact that this server is <i>still</i> down or unreachable. Unless you've changed the <a href="#">interval_length</a> directive from the default value of 60, this number will mean minutes. If you set this value to 0, Nagios will <i>not</i> re-notify contacts about problems for this host - only one problem notification will be sent out.                                                                                                                                                                                                                 |
| <b>notification_period:</b>     | This directive is used to specify the short name of the <a href="#">time period</a> during which notifications of events for this host can be sent out to contacts. If a host goes down, becomes unreachable, or recovers during a time which is not covered by the time period, no notifications will be sent out.                                                                                                                                                                                                                                                                                                                                           |
| <b>notification_options:</b>    | This directive is used to determine when notifications for the host should be sent out. Valid options are a combination of one or more of the following: <b>d</b> = send notifications on a DOWN state, <b>u</b> = send notifications on an UNREACHABLE state, <b>r</b> = send notifications on recoveries (OK state), and <b>f</b> = send notifications when the host starts and stops <a href="#">flapping</a> . If you specify <b>n</b> (none) as an option, no host notifications will be sent out. Example: If you specify <b>d,r</b> in this field, notifications will only be sent out when the host goes DOWN and when it recovers from a DOWN state. |
| <b>notifications_enabled</b> *: | This directive is used to determine whether or not notifications for this host are enabled. Values: 0 = disable host notifications, 1 = enable host notifications.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
| <b>stalking_options:</b>        | This directive determines which host states "stalking" is enabled for. Valid options are a combination of one or more of the following: <b>o</b> = stalk on UP states, <b>d</b> = stalk on DOWN states, and <b>u</b> = stalk on UNREACHABLE states. More information on state stalking can be found <a href="#">here</a> .                                                                                                                                                                                                                                                                                                                                    |

## Host Group Definition

### Description:

A host group definition is used to group one or more hosts together for display purposes in the [CGIs](#).

### Definition Format:

Note: Directives in red are required, while those in black are optional.

```
define hostgroup{
    hostgroup_name hostgroup_name
    alias          alias
    members        members
}
```

### Example Definition:

```
define hostgroup{
    hostgroup_name    novell-servers
    alias              Novell Servers
    members            netware1,netware2,netware3,netware4
}
```

### Directive Descriptions:

- hostgroup\_name:** This directive is used to define a short name used to identify the host group.
- alias:** This directive is used to define is a longer name or description used to identify the host group. It is provided in order to allow you to more easily identify a particular host group.
- members:** This is a list of the *short names* of [hosts](#) that should be included in this group. Multiple host names should be separated by commas. This directive may be used as an alternative to (or in addition to) the *hostgroups* directive in [host definitions](#).

### Service Definition

#### Description:

A service definition is used to identify a "service" that runs on a host. The term "service" is used very loosely. It can mean an actual service that runs on the host (POP, SMTP, HTTP, etc.) or some other type of metric associated with the host (response to a ping, number of logged in users, free disk space, etc.). The different arguments to a service definition are outlined below.

#### Definition Format:

Note: Directives in red are required, while those in black are optional.

```

define service{
    host_name                host_name
    service_description       service_description
    servicegroups             servicegroup_names
    is_volatile               [0/1]
    check_command              command_name
    max_check_attempts        #
    normal_check_interval     #
    retry_check_interval      #
    active_checks_enabled     [0/1]
    passive_checks_enabled    [0/1]
    check_period              timeperiod_name
    parallelize_check         [0/1]
    obsess_over_service       [0/1]
    check_freshness           [0/1]
    freshness_threshold       #
    event_handler              command_name
    event_handler_enabled     [0/1]
    low_flap_threshold        #
    high_flap_threshold       #
    flap_detection_enabled    [0/1]
    process_perf_data         [0/1]
    retain_status_information  [0/1]
    retain_nonstatus_information [0/1]
    notification_interval     #
    notification_period       timeperiod_name
    notification_options       [w,u,c,r,f]
    notifications_enabled      [0/1]
    contact_groups             contact_groups
    stalking_options           [o,w,u,c]
}

```

## Example Definition:

```
define service{
    host_name                linux-server
    service_description      check-disk-sda1
    check_command             check-disk!/dev/sda1
    max_check_attempts       5
    normal_check_interval    5
    retry_check_interval     3
    check_period             24x7
    notification_interval    30
    notification_period      24x7
    notification_options     w,c,r
    contact_groups           linux-admins
}
```

## Directive Descriptions:

|                               |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
|-------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>host_name:</b>             | This directive is used to specify the <i>short name</i> of the <a href="#">host</a> that the service "runs" on or is associated with.                                                                                                                                                                                                                                                                                                                                                                                                                     |
| <b>service_description:</b>   | This directive is used to define the description of the service, which may contain spaces, dashes, and colons (semicolons, apostrophes, and quotation marks should be avoided). No two services associated with the same host can have the same description. Services are uniquely identified with their <i>host_name</i> and <i>service_description</i> directives.                                                                                                                                                                                      |
| <b>servicegroups:</b>         | This directive is used to identify the <i>short name(s)</i> of the <a href="#">servicegroup(s)</a> that the service belongs to. Multiple servicegroups should be separated by commas. This directive may be used as an alternative to using the <i>members</i> directive in <a href="#">servicegroup</a> definitions.                                                                                                                                                                                                                                     |
| <b>is_volatile:</b>           | This directive is used to denote whether the service is "volatile". Services are normally <i>not</i> volatile. More information on volatile service and how they differ from normal services can be found <a href="#">here</a> . Value: 0 = service is not volatile, 1 = service is volatile.                                                                                                                                                                                                                                                             |
| <b>check_command:</b>         | This directive is used to specify the <i>short name</i> of the <a href="#">command</a> that Nagios will run in order to check the status of the service. The maximum amount of time that the service check command can run is controlled by the <a href="#">service_check_timeout</a> option.                                                                                                                                                                                                                                                             |
| <b>max_check_attempts:</b>    | This directive is used to define the number of times that Nagios will retry the service check command if it returns any state other than an OK state. Setting this value to 1 will cause Nagios to generate an alert without retrying the service check again.                                                                                                                                                                                                                                                                                            |
| <b>normal_check_interval:</b> | This directive is used to define the number of "time units" to wait before scheduling the next "regular" check of the service. "Regular" checks are those that occur when the service is in an OK state or when the service is in a non-OK state, but has already been rechecked <b>max_attempts</b> number of times. Unless you've changed the <a href="#">interval_length</a> directive from the default value of 60, this number will mean minutes. More information on this value can be found in the <a href="#">check scheduling</a> documentation. |



|                                  |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 |
|----------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>retry_check_interval:</b>     | This directive is used to define the number of "time units" to wait before scheduling a re-check of the service. Services are rescheduled at the retry interval when they have changed to a non-OK state. Once the service has been retried <b>max_attempts</b> times without a change in its status, it will revert to being scheduled at its "normal" rate as defined by the <b>check_interval</b> value. Unless you've changed the <a href="#">interval_length</a> directive from the default value of 60, this number will mean minutes. More information on this value can be found in the <a href="#">check scheduling</a> documentation. |
| <b>active_checks_enabled *:</b>  | This directive is used to determine whether or not active checks of this service are enabled. Values: 0 = disable active service checks, 1 = enable active service checks.                                                                                                                                                                                                                                                                                                                                                                                                                                                                      |
| <b>passive_checks_enabled *:</b> | This directive is used to determine whether or not passive checks of this service are enabled. Values: 0 = disable passive service checks, 1 = enable passive service checks.                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
| <b>check_period:</b>             | This directive is used to specify the short name of the <a href="#">time period</a> during which active checks of this service can be made.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
| <b>parallelize_check:</b>        | This directive is used to determine whether or not the service check can be parallelized. By default, all service checks are parallelized. Disabling parallel checks of services can result in serious performance problems. More information on service check parallelization can be found <a href="#">here</a> . Values: 0 = service check cannot be parallelized (use with caution!), 1 = service check can be parallelized.                                                                                                                                                                                                                 |
| <b>obsess_over_service *:</b>    | This directive determines whether or not checks for the service will be "obsessed" over using the <a href="#">ocsp_command</a> .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| <b>check_freshness *:</b>        | This directive is used to determine whether or not <a href="#">freshness checks</a> are enabled for this service. Values: 0 = disable freshness checks, 1 = enable freshness checks.                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
| <b>freshness_threshold:</b>      | This directive is used to specify the freshness threshold (in seconds) for this service. If you set this directive to a value of 0, Nagios will determine a freshness threshold to use automatically.                                                                                                                                                                                                                                                                                                                                                                                                                                           |
| <b>event_handler_enabled *:</b>  | This directive is used to determine whether or not the event handler for this service is enabled. Values: 0 = disable service event handler, 1 = enable service event handler.                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |
| <b>low_flap_threshold:</b>       | This directive is used to specify the low state change threshold used in flap detection for this service. More information on flap detection can be found <a href="#">here</a> . If you set this directive to a value of 0, the program-wide value specified by the <a href="#">low_service_flap_threshold</a> directive will be used.                                                                                                                                                                                                                                                                                                          |
| <b>high_flap_threshold:</b>      | This directive is used to specify the high state change threshold used in flap detection for this service. More information on flap detection can be found <a href="#">here</a> . If you set this directive to a value of 0, the program-wide value specified by the <a href="#">high_service_flap_threshold</a> directive will be used.                                                                                                                                                                                                                                                                                                        |

|                                      |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
|--------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>flap_detection_enabled *:</b>     | This directive is used to determine whether or not flap detection is enabled for this service. More information on flap detection can be found <a href="#">here</a> . Values: 0 = disable service flap detection, 1 = enable service flap detection.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
| <b>process_perf_data *:</b>          | This directive is used to determine whether or not the processing of performance data is enabled for this service. Values: 0 = disable performance data processing, 1 = enable performance data processing.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    |
| <b>retain_status_information:</b>    | This directive is used to determine whether or not status-related information about the service is retained across program restarts. This is only useful if you have enabled state retention using the <a href="#">retain_state_information</a> directive. Value: 0 = disable status information retention, 1 = enable status information retention.                                                                                                                                                                                                                                                                                                                                                                                           |
| <b>retain_nonstatus_information:</b> | This directive is used to determine whether or not non-status information about the service is retained across program restarts. This is only useful if you have enabled state retention using the <a href="#">retain_state_information</a> directive. Value: 0 = disable non-status information retention, 1 = enable non-status information retention.                                                                                                                                                                                                                                                                                                                                                                                       |
| <b>notification_interval:</b>        | This directive is used to define the number of "time units" to wait before re-notifying a contact that this service is <i>still</i> in a non-OK state. Unless you've changed the <a href="#">interval_length</a> directive from the default value of 60, this number will mean minutes. If you set this value to 0, Nagios will <i>not</i> re-notify contacts about problems for this service - only one problem notification will be sent out.                                                                                                                                                                                                                                                                                                |
| <b>notification_period:</b>          | This directive is used to specify the short name of the <a href="#">time period</a> during which notifications of events for this service can be sent out to contacts. No service notifications will be sent out during times which is not covered by the time period.                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
| <b>notification_options:</b>         | This directive is used to determine when notifications for the service should be sent out. Valid options are a combination of one or more of the following: <b>w</b> = send notifications on a WARNING state, <b>u</b> = send notifications on an UNKNOWN state, <b>c</b> = send notifications on a CRITICAL state, <b>r</b> = send notifications on recoveries (OK state), and <b>f</b> = send notifications when the service starts and stops <a href="#">flapping</a> . If you specify <b>n</b> (none) as an option, no service notifications will be sent out. Example: If you specify <b>w,r</b> in this field, notifications will only be sent out when the service goes into a WARNING state and when it recovers from a WARNING state. |
| <b>notifications_enabled *:</b>      | This directive is used to determine whether or not notifications for this service are enabled. Values: 0 = disable service notifications, 1 = enable service notifications.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    |
| <b>contact_groups:</b>               | This is a list of the <i>short names</i> of the <a href="#">contact groups</a> that should be notified whenever there are problems (or recoveries) with this service. Multiple contact groups should be separated by commas.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |

**stalking\_options:** This directive determines which service states "stalking" is enabled for. Valid options are a combination of one or more of the following: **o** = stalk on OK states, **w** = stalk on WARNING states, **u** = stalk on UNKNOWN states, and **c** = stalk on CRITICAL states. More information on state stalking can be found [here](#).

## Service Group Definition

### Description:

A service group definition is used to group one or more services together for display purposes in the [CGIs](#).

### Definition Format:

Note: Directives in red are required, while those in black are optional.

```
define servicegroup{
    servicegroup_name  servicegroup_name
    alias               alias
    members             members
}
```

### Example Definition:

```
define servicegroup{
    servicegroup_name  dbservices
    alias              Database Services
    members             ms1,SQL Server,ms1,SQL Server Agent,ms1,SQL DTC
}
```

### Directive Descriptions:

**servicegroup\_name:** This directive is used to define a short name used to identify the service group.

**alias:** This directive is used to define is a longer name or description used to identify the service group. It is provided in order to allow you to more easily identify a particular service group.

**members:** This is a list of the *descriptions* of [services](#) (and the names of their corresponding hosts) that should be included in this group. Host and service names should be separated by commas. This directive may be used as an alternative to the *servicegroups* directive in [service definitions](#). The format of the member directive is as follows (note that a host name must precede a service name/description):

members=<host1>,<service1>,<host2>,<service2>,...,<hostn>,<servicen>

## Contact Definition

### Description:

A contact definition is used to identify someone who should be contacted in the event of a problem on your network. The different arguments to a contact definition are described below.

### Definition Format:

Note: Directives in red are required, while those in black are optional.

```
define contact{  
    contact_name          contact_name  
    alias                  alias  
    contactgroups          contactgroup_names  
    host_notification_period timeperiod_name  
    service_notification_period timeperiod_name  
    host_notification_options [d,u,r,f,n]  
    service_notification_options [w,u,c,r,f,n]  
    host_notification_commands command_name  
    service_notification_commands command_name  
    email                  email_address  
    pager                  pager_number or pager_email_gateway  
    addressx               additional_contact_address  
}
```

### Example Definition:

```
define contact{  
    contact_name          jdoe  
    alias                  John Doe  
    service_notification_period 24x7  
    host_notification_period 24x7  
    service_notification_options w,u,c,r  
    host_notification_options d,u,r  
    service_notification_commands notify-by-email  
    host_notification_commands host-notify-by-email  
    email                  jdoe@localhost.localdomain  
    pager                  555-5555@pagergateway.localhost.localdomain  
    address1               xxxxx.xyyy@icq.com  
    address2               555-555-5555  
}
```

### Directive Descriptions:

|                                      |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
|--------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>contact_name:</b>                 | This directive is used to define a short name used to identify the contact. It is referenced in <a href="#">contact group</a> definitions. Under the right circumstances, the <code>\$CONTACTNAME\$</code> <a href="#">macro</a> will contain this value.                                                                                                                                                                                                                                                                                                                  |
| <b>alias:</b>                        | This directive is used to define a longer name or description for the contact. Under the rights circumstances, the <code>\$CONTACTALIAS\$</code> <a href="#">macro</a> will contain this value.                                                                                                                                                                                                                                                                                                                                                                            |
| <b>contactgroups:</b>                | This directive is used to identify the <i>short name(s)</i> of the <a href="#">contactgroup(s)</a> that the contact belongs to. Multiple <a href="#">contactgroups</a> should be separated by commas. This directive may be used as an alternative to (or in addition to) using the <i>members</i> directive in <a href="#">contactgroup</a> definitions.                                                                                                                                                                                                                  |
| <b>host_notification_period:</b>     | This directive is used to specify the short name of the <a href="#">time period</a> during which the contact can be notified about host problems or recoveries. You can think of this as an "on call" time for host notifications for the contact. Read the documentation on <a href="#">time periods</a> for more information on how this works and potential problems that may result from improper use.                                                                                                                                                                 |
| <b>service_notification_period:</b>  | This directive is used to specify the short name of the <a href="#">time period</a> during which the contact can be notified about service problems or recoveries. You can think of this as an "on call" time for service notifications for the contact. Read the documentation on <a href="#">time periods</a> for more information on how this works and potential problems that may result from improper use.                                                                                                                                                           |
| <b>host_notification_commands:</b>   | This directive is used to define a list of the <i>short names</i> of the <a href="#">commands</a> used to notify the contact of a <i>host</i> problem or recovery. Multiple notification commands should be separated by commas. All notification commands are executed when the contact needs to be notified. The maximum amount of time that a notification command can run is controlled by the <a href="#">notification_timeout</a> option.                                                                                                                            |
| <b>host_notification_options:</b>    | This directive is used to define the host states for which notifications can be sent out to this contact. Valid options are a combination of one or more of the following: <b>d</b> = notify on DOWN host states, <b>u</b> = notify on UNREACHABLE host states, <b>r</b> = notify on host recoveries (UP states), and <b>f</b> = notify when the host starts and stops <a href="#">flapping</a> . If you specify <b>n</b> (none) as an option, the contact will not receive any type of host notifications.                                                                |
| <b>service_notification_options:</b> | This directive is used to define the service states for which notifications can be sent out to this contact. Valid options are a combination of one or more of the following: <b>w</b> = notify on WARNING service states, <b>u</b> = notify on UNKNOWN service states, <b>c</b> = notify on CRITICAL service states, <b>r</b> = notify on service recoveries (OK states), and <b>f</b> = notify when the service starts and stops <a href="#">flapping</a> . If you specify <b>n</b> (none) as an option, the contact will not receive any type of service notifications. |

|                                       |                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
|---------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>service_notification_commands:</b> | This directive is used to define a list of the <i>short names</i> of the <b>commands</b> used to notify the contact of a <i>service</i> problem or recovery. Multiple notification commands should be separated by commas. All notification commands are executed when the contact needs to be notified. The maximum amount of time that a notification command can run is controlled by the <b>notification_timeout</b> option.                           |
| <b>email:</b>                         | This directive is used to define an email address for the contact. Depending on how you configure your notification commands, it can be used to send out an alert email to the contact. Under the right circumstances, the \$CONTACTEMAIL\$ <b>macro</b> will contain this value.                                                                                                                                                                          |
| <b>pager:</b>                         | This directive is used to define a pager number for the contact. It can also be an email address to a pager gateway (i.e. pagejoe@pagenet.com). Depending on how you configure your notification commands, it can be used to send out an alert page to the contact. Under the right circumstances, the \$CONTACTPAGER\$ <b>macro</b> will contain this value.                                                                                              |
| <b>addressx:</b>                      | Address directives are used to define additional "addresses" for the contact. These addresses can be anything - cell phone numbers, instant messaging addresses, etc. Depending on how you configure your notification commands, they can be used to send out an alert o the contact. Up to six addresses can be defined using these directives ( <i>address1</i> through <i>address6</i> ). The \$CONTACTADDRESSx\$ <b>macro</b> will contain this value. |

## Contact Group Definition

### Description:

A contact group definition is used to group one or more **contacts** together for the purpose of sending out alert/recovery notifications. When a host or service has a problem or recovers, Nagios will find the appropriate contact groups to send notifications to, and notify all contacts in those contact groups. This may sound complex, but for most people it doesn't have to be. It does, however, allow for flexibility in determining who gets notified for particular events. The different arguments to a contact group definition are outlined below.

### Definition Format:

Note: Directives in red are required, while those in black are optional.

```
define contactgroup{
    contactgroup_name    contactgroup_name
    alias                alias
    members              members
}
```

### Example Definition:

```
define contactgroup{
    contactgroup_name    novell-admins
    alias                Novell Administrators
    members              jdoe,rtobert,tzach
}
```

### Directive Descriptions:

- contactgroup\_name:** This directive is a short name used to identify the contact group.
- alias:** This directive is used to define a longer name or description used to identify the contact group.
- members:** This directive is used to define a list of the *short names* of [contacts](#) that should be included in this group. Multiple contact names should be separated by commas. This directive may be used as an alternative to (or in addition to) using the *contactgroups* directive in [contact](#) definitions.

### Time Period Definition

#### Description:

A time period is a list of times during various days that are considered to be "valid" times for notifications and service checks. It consists one or more time periods for each day of the week that "rotate" once the week has come to an end. Exceptions to the normal weekly time range rotations are not supported.

#### Definition Format:

Note: Directives in red are required, while those in black are optional.

```
define timeperiod{
    timeperiod_name    timeperiod_name
    alias              alias
    sunday             timeranges
    monday             timeranges
    tuesday            timeranges
    wednesday          timeranges
    thursday           timeranges
    friday             timeranges
    saturday           timeranges
}
```

## Example Definition:

```
define timeperiod{
    timeperiod_name    nonworkhours
    alias               Non-Work Hours
    sunday              00:00-24:00
    monday              00:00-09:00,17:00-24:00
    tuesday             00:00-09:00,17:00-24:00
    wednesday           00:00-09:00,17:00-24:00
    thursday            00:00-09:00,17:00-24:00
    friday              00:00-09:00,17:00-24:00
    saturday            00:00-24:00
}
```

## Directive Descriptions:

- timeperiod\_name:** This directive is the short name used to identify the time period.
- alias:** This directive is a longer name or description used to identify the time period.
- someday:** The *sunday* through *saturday* directives are comma-delimited lists of time ranges that are "valid" times for a particular day of the week. Notice that there are seven different days for which you can define time ranges (Sunday through Saturday). Each time range is in the form of **HH:MM-HH:MM**, where hours are specified on a 24 hour clock. For example, **00:15-24:00** means 12:15am in the morning for this day until 12:20am midnight (a 23 hour, 45 minute total time range). If you wish to exclude an entire day from the timeperiod, simply do not include it in the timeperiod definition.

## Command Definition

### Description:

A command definition is just that. It defines a command. Commands that can be defined include service checks, service notifications, service event handlers, host checks, host notifications, and host event handlers. Command definitions can contain [macros](#), but you must make sure that you include only those macros that are "valid" for the circumstances when the command will be used. More information on what macros are available and when they are "valid" can be found [here](#). The different arguments to a command definition are outlined below.

### Definition Format:

Note: Directives in red are required, while those in black are optional.

```
define command{
    command_name    command_name
    command_line    command_line
}
```

## Example Definition:



```
define command{
    command_name    check_pop
    command_line    /usr/local/nagios/libexec/check_pop -H $HOSTADDRESS$
}
```

## Directive Descriptions:

**command\_name:** This directive is the short name used to identify the command. It is referenced in [contact](#), [host](#), and [service](#) definitions (in notification, check, and event handler directives), among other places.

**command\_line:** This directive is used to define what is actually executed by Nagios when the command is used for service or host checks, notifications, or [event handlers](#). Before the command line is executed, all valid [macros](#) are replaced with their respective values. See the documentation on macros for determining when you can use different macros. Note that the command line is *not* surrounded in quotes. Also, if you want to pass a dollar sign (\$) on the command line, you have to escape it with another dollar sign.

**NOTE:** You may not include a **semicolon** (;) in the *command\_line* directive, because everything after it will be ignored as a config file comment. You can work around this limitation by setting one of the [\\$USER\\$ macros in your resource file](#) to a semicolon and then referencing the appropriate \$USER\$ macro in the *command\_line* directive in place of the semicolon.

If you want to pass arguments to commands during runtime, you can use [\\$ARGn\\$ macros](#) in the *command\_line* directive of the command definition and then separate individual arguments from the command name (and from each other) using bang (!) characters in the object definition directive (host check command, service event handler command, etc) that references the command. More information on how arguments in command definitions are processed during runtime can be found in the documentation on [macros](#).

## Service Dependency Definition

### Description:

Service dependencies are an advanced feature of Nagios that allow you to suppress notifications and active checks of services based on the status of one or more other services. Service dependencies are optional and are mainly targeted at advanced users who have complicated monitoring setups. More information on how service dependencies work (read this!) can be found [here](#).

### Definition Format:

Note: Directives in red are required, while those in black are optional. However, you must supply at least one type of criteria for the definition to be of much use.

```

define servicedependency{
    dependent_host_name      host_name
    dependent_service_description  service_description
    host_name                host_name
    service_description      service_description
    inherits_parent          [0/1]
    execution_failure_criteria  [o,w,u,c,p,n]
    notification_failure_criteria  [o,w,u,c,p,n]
}

```

Example Definition:

```

define servicedependency{
    host_name                WWW1
    service_description      Apache Web Server
    dependent_host_name      WWW1
    dependent_service_description  Main Web Site
    execution_failure_criteria  n
    notification_failure_criteria  w,u,c
}

```

Directive Descriptions:

|                                       |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
|---------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>dependent_host:</b>                | This directive is used to identify the <i>short name</i> of the <b>host</b> that the <i>dependent</i> service "runs" on or is associated with.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |
| <b>dependent_service_description:</b> | This directive is used to identify the <i>description</i> of the <i>dependent service</i> .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |
| <b>host_name:</b>                     | This directive is used to identify the <i>short name</i> of the <b>host</b> that the service <i>that is being depended upon</i> (also referred to as the master service) "runs" on or is associated with.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| <b>service_description:</b>           | This directive is used to identify the <i>description</i> of the <b>service</b> <i>that is being depended upon</i> (also referred to as the master service).                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
| <b>inherits_parent:</b>               | This directive indicates whether or not the dependency inherits dependencies of the service <i>that is being depended upon</i> (also referred to as the master service). In other words, if the master service is dependent upon other services and any one of those dependencies fail, this dependency will also fail.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 |
| <b>execution_failure_criteria:</b>    | This directive is used to specify the criteria that determine when the dependent service should <i>not</i> be actively checked. If the <i>master</i> service is in one of the failure states we specify, the <i>dependent</i> service will not be actively checked. Valid options are a combination of one or more of the following (multiple options are separated with commas): <b>o</b> = fail on an OK state, <b>w</b> = fail on a WARNING state, <b>u</b> = fail on an UNKNOWN state, <b>c</b> = fail on a CRITICAL state, and <b>p</b> = fail on a pending state (e.g. the service has not yet been checked). If you specify <b>n</b> (none) as an option, the execution dependency will never fail and checks of the dependent service will always be actively checked (if other conditions allow for it to be). Example: If you specify <b>o,c,u</b> in this field, the <i>dependent</i> service will not be actively checked if the <i>master</i> service is in either an OK, a CRITICAL, or an UNKNOWN state. |
| <b>notification_failure_criteria:</b> | This directive is used to define the criteria that determine when notifications for the dependent service should <i>not</i> be sent out. If the <i>master</i> service is in one of the failure states we specify, notifications for the <i>dependent</i> service will not be sent to contacts. Valid options are a combination of one or more of the following: <b>o</b> = fail on an OK state, <b>w</b> = fail on a WARNING state, <b>u</b> = fail on an UNKNOWN state, <b>c</b> = fail on a CRITICAL state, and <b>p</b> = fail on a pending state (e.g. the service has not yet been checked). If you specify <b>n</b> (none) as an option, the notification dependency will never fail and notifications for the dependent service will always be sent out. Example: If you specify <b>w</b> in this field, the notifications for the <i>dependent</i> service will not be sent out if the <i>master</i> service is in a WARNING state.                                                                             |

## Service Escalation Definition

Description:

Service escalations are *completely optional* and are used to escalate notifications for a particular service. More information on how notification escalations work can be found [here](#).

Definition Format:

Note: Directives in red are required, while those in black are optional.

```
define serviceescalation{
    host_name            host_name
    service_description  service_description
    contact_groups        contactgroup_name
    first_notification    #
    last_notification     #
    notification_interval #
    escalation_period     timeperiod_name
    escalation_options    [w,u,c,r]
}
```

Example Definition:

```
define serviceescalation{
    host_name            nt-3
    service_description  Processor Load
    first_notification    4
    last_notification     0
    notification_interval 30
    contact_groups        all-nt-admins,themanagers
}
```

Directive Descriptions:

|                               |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
|-------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>host_name:</b>             | This directive is used to identify the <i>short name</i> of the <a href="#">host</a> that the <a href="#">service</a> the escalation should apply to is associated with.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| <b>service_description:</b>   | This directive is used to identify the <i>description</i> of the <a href="#">service</a> the escalation should apply to.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| <b>first_notification:</b>    | This directive is a number that identifies the <i>first</i> notification for which this escalation is effective. For instance, if you set this value to 3, this escalation will only be used if the service is in a non-OK state long enough for a third notification to go out.                                                                                                                                                                                                                                                                                                                                                                                                                        |
| <b>last_notification:</b>     | This directive is a number that identifies the <i>last</i> notification for which this escalation is effective. For instance, if you set this value to 5, this escalation will not be used if more than five notifications are sent out for the service. Setting this value to 0 means to keep using this escalation entry forever (no matter how many notifications go out).                                                                                                                                                                                                                                                                                                                           |
| <b>contact_groups:</b>        | This directive is used to identify the <i>short name</i> of the <a href="#">contact group</a> that should be notified when the service notification is escalated. Multiple contact groups should be separated by commas.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| <b>notification_interval:</b> | This directive is used to determine the interval at which notifications should be made while this escalation is valid. If you specify a value of 0 for the interval, Nagios will send the first notification when this escalation definition is valid, but will then prevent any more problem notifications from being sent out for the host. Notifications are sent out again until the host recovers. This is useful if you want to stop having notifications sent out after a certain amount of time. Note: If multiple escalation entries for a host overlap for one or more notification ranges, the smallest notification interval from all escalation entries is used.                           |
| <b>escalation_period:</b>     | This directive is used to specify the short name of the <a href="#">time period</a> during which this escalation is valid. If this directive is not specified, the escalation is considered to be valid during all times.                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| <b>escalation_options:</b>    | This directive is used to define the criteria that determine when this service escalation is used. The escalation is used only if the service is in one of the states specified in this directive. If this directive is not specified in a service escalation, the escalation is considered to be valid during all service states. Valid options are a combination of one or more of the following: <b>r</b> = escalate on an OK (recovery) state, <b>w</b> = escalate on a WARNING state, <b>u</b> = escalate on an UNKNOWN state, and <b>c</b> = escalate on a CRITICAL state. Example: If you specify <b>w</b> in this field, the escalation will only be used if the service is in a WARNING state. |

## Host Dependency Definition

### Description:

Host dependencies are an advanced feature of Nagios that allow you to suppress notifications for hosts based on the status of one or more other hosts. Host dependencies are optional and are mainly targeted at advanced users who have complicated monitoring setups. More information on how host dependencies work (read this!) can be found [here](#).

## Definition Format:

Note: Directives in red are required, while those in black are optional.

```
define hostdependency{  
    dependent_host_name      host_name  
    host_name                host_name  
    inherits_parent          [0/1]  
    execution_failure_criteria [o,d,u,p,n]  
    notification_failure_criteria [o,d,u,p,n]  
}
```

## Example Definition:

```
define hostdependency{  
    host_name                WWW1  
    dependent_host_name      DBASE1  
    notification_failure_criteria d,u  
}
```

## Directive Descriptions:

|                                       |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |
|---------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>dependent_host:</b>                | This directive is used to identify the <i>short name</i> of the <i>dependent</i> <a href="#">host</a> .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |
| <b>host_name:</b>                     | This directive is used to identify the <i>short name</i> of the <a href="#">host</a> <i>that is being depended upon</i> (also referred to as the master host).                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
| <b>inherits_parent:</b>               | This directive indicates whether or not the dependency inherits dependencies of the host <i>that is being depended upon</i> (also referred to as the master host). In other words, if the master host is dependent upon other hosts and any one of those dependencies fail, this dependency will also fail.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
| <b>execution_failure_criteria:</b>    | This directive is used to specify the criteria that determine when the dependent host should <i>not</i> be actively checked. If the <i>master</i> host is in one of the failure states we specify, the <i>dependent</i> host will not be actively checked. Valid options are a combination of one or more of the following (multiple options are separated with commas): <b>o</b> = fail on an UP state, <b>d</b> = fail on a DOWN state, <b>u</b> = fail on an UNREACHABLE state, and <b>p</b> = fail on a pending state (e.g. the host has not yet been checked). If you specify <b>n</b> (none) as an option, the execution dependency will never fail and the dependent host will always be actively checked (if other conditions allow for it to be). Example: If you specify <b>u,d</b> in this field, the <i>dependent</i> host will not be actively checked if the <i>master</i> host is in either an UNREACHABLE or DOWN state. |
| <b>notification_failure_criteria:</b> | This directive is used to define the criteria that determine when notifications for the dependent host should <i>not</i> be sent out. If the <i>master</i> host is in one of the failure states we specify, notifications for the <i>dependent</i> host will not be sent to contacts. Valid options are a combination of one or more of the following: <b>o</b> = fail on an UP state, <b>d</b> = fail on a DOWN state, <b>u</b> = fail on an UNREACHABLE state, and <b>p</b> = fail on a pending state (e.g. the host has not yet been checked). If you specify <b>n</b> (none) as an option, the notification dependency will never fail and notifications for the dependent host will always be sent out. Example: If you specify <b>d</b> in this field, the notifications for the <i>dependent</i> host will not be sent out if the <i>master</i> host is in a DOWN state.                                                          |

## Host Escalation Definition

### Description:

Host escalations are *completely optional* and are used to escalate notifications for a particular host. More information on how notification escalations work can be found [here](#).

### Definition Format:

Note: Directives in red are required, while those in black are optional.

```

define hostescalation{
    host_name          host_name
    hostgroup_name     hostgroup_name
    contact_groups     contactgroup_name
    first_notification  #
    last_notification   #
    notification_interval #
    escalation_period   timeperiod_name
    escalation_options  [d,u,r]
}

```

#### Example Definition:

```

define hostescalation{
    host_name          router-34
    first_notification  5
    last_notification   8
    notification_interval 60
    contact_groups     all-router-admins
}

```

#### Directive Descriptions:



|                               |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|-------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>host_name:</b>             | This directive is used to identify the <i>short name</i> of the <a href="#">host</a> that the escalation should apply to.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
| <b>hostgroup_name:</b>        | This directive is used to identify the <i>short name(s)</i> of the <a href="#">hostgroup(s)</a> that the escalation should apply to. Multiple hostgroups should be separated by commas. If this is used, the escalation will apply to all hosts that are members of the specified hostgroup(s).                                                                                                                                                                                                                                                                                                                                                                               |
| <b>first_notification:</b>    | This directive is a number that identifies the <i>first</i> notification for which this escalation is effective. For instance, if you set this value to 3, this escalation will only be used if the host is down or unreachable long enough for a third notification to go out.                                                                                                                                                                                                                                                                                                                                                                                               |
| <b>last_notification:</b>     | This directive is a number that identifies the <i>last</i> notification for which this escalation is effective. For instance, if you set this value to 5, this escalation will not be used if more than five notifications are sent out for the host. Setting this value to 0 means to keep using this escalation entry forever (no matter how many notifications go out).                                                                                                                                                                                                                                                                                                    |
| <b>contact_groups:</b>        | This directive is used to identify the <i>short name</i> of the <a href="#">contact group</a> that should be notified when the host notification is escalated. Multiple contact groups should be separated by commas.                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
| <b>notification_interval:</b> | This directive is used to determine the interval at which notifications should be made while this escalation is valid. If you specify a value of 0 for the interval, Nagios will send the first notification when this escalation definition is valid, but will then prevent any more problem notifications from being sent out for the host. Notifications are sent out again until the host recovers. This is useful if you want to stop having notifications sent out after a certain amount of time. Note: If multiple escalation entries for a host overlap for one or more notification ranges, the smallest notification interval from all escalation entries is used. |
| <b>escalation_period:</b>     | This directive is used to specify the short name of the <a href="#">time period</a> during which this escalation is valid. If this directive is not specified, the escalation is considered to be valid during all times.                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
| <b>escalation_options:</b>    | This directive is used to define the criteria that determine when this host escalation is used. The escalation is used only if the host is in one of the states specified in this directive. If this directive is not specified in a host escalation, the escalation is considered to be valid during all host states. Valid options are a combination of one or more of the following: <b>r</b> = escalate on an UP (recovery) state, <b>d</b> = escalate on a DOWN state, and <b>u</b> = escalate on an UNREACHABLE state. Example: If you specify <b>d</b> in this field, the escalation will only be used if the host is in a DOWN state.                                 |

## Extended Host Information Definition

### Description:

Extended host information entries are basically used to make the output from the [status](#), [statusmap](#), [statuswrl](#), and [extinfo](#) CGIs look pretty. They have no effect on monitoring and are completely optional.

## Definition Format:

Note: Variables in red are required, while those in black are optional. However, you need to supply at least one optional variable in each definition for it to be of much use.

```
define hostextinfo{  
    host_name      host_name  
    notes          note_string  
    notes_url      url  
    action_url     url  
    icon_image     image_file  
    icon_image_alt alt_string  
    vrml_image     image_file  
    statusmap_image image_file  
    2d_coords      x_coord,y_coord  
    3d_coords      x_coord,y_coord,z_coord  
}
```

## Example Definition:

```
define hostextinfo{  
    host_name      network1  
    notes          This is the primary Netware file server  
    notes_url      http://webserver.localhost.localdomain/hostinfo.pl?host=network1  
    icon_image     novell40.png  
    icon_image_alt IntranetWare 4.11  
    vrml_image     novell40.png  
    statusmap_image novell40.gd2  
    2d_coords      100,250  
    3d_coords      100.0,50.0,75.0  
}
```

## Variable Descriptions:

- host\_name:** This variable is used to identify the *short name* of the [host](#) which the data is associated with.
- notes:** This directive is used to define an optional string of notes pertaining to the host. If you specify a note here, you will see the it in the [extended information](#) CGI (when you are viewing information about the specified host).

|                         |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |
|-------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>notes_url:</b>       | This variable is used to define an optional URL that can be used to provide more information about the host. If you specify an URL, you will see a link that says "Extra Host Notes" in the <a href="#">extended information</a> CGI (when you are viewing information about the specified host). Any valid URL can be used. If you plan on using relative paths, the base path will be the same as what is used to access the CGIs (i.e. <code>/cgi-bin/nagios/</code> ). This can be very useful if you want to make detailed information on the host, emergency contact methods, etc. available to other support staff.                                                                                                                                                                                                                                                                                                                                       |
| <b>action_url:</b>      | This directive is used to define an optional URL that can be used to provide more actions to be performed on the host. If you specify an URL, you will see a link that says "Extra Host Actions" in the <a href="#">extended information</a> CGI (when you are viewing information about the specified host). Any valid URL can be used. If you plan on using relative paths, the base path will be the same as what is used to access the CGIs (i.e. <code>/cgi-bin/nagios/</code> ).                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
| <b>icon_image:</b>      | This variable is used to define the name of a GIF, PNG, or JPG image that should be associated with this host. This image will be displayed in the <a href="#">status</a> and <a href="#">extended information</a> CGIs. The image will look best if it is 40x40 pixels in size. Images for hosts are assumed to be in the <b>logos/</b> subdirectory in your HTML images directory (i.e. <code>/usr/local/nagios/share/images/logos</code> ).                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
| <b>icon_image_alt:</b>  | This variable is used to define an optional string that is used in the ALT tag of the image specified by the <code>&lt;icon_image&gt;</code> argument. The ALT tag is used in the <a href="#">status</a> , <a href="#">extended information</a> and <a href="#">statusmap</a> CGIs.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
| <b>vrml_image:</b>      | This variable is used to define the name of a GIF, PNG, or JPG image that should be associated with this host. This image will be used as the texture map for the specified host in the <a href="#">statuswrl</a> CGI. Unlike the image you use for the <code>&lt;icon_image&gt;</code> variable, this one should probably <i>not</i> have any transparency. If it does, the host object will look a bit wierd. Images for hosts are assumed to be in the <b>logos/</b> subdirectory in your HTML images directory (i.e. <code>/usr/local/nagios/share/images/logos</code> ).                                                                                                                                                                                                                                                                                                                                                                                    |
| <b>statusmap_image:</b> | This variable is used to define the name of an image that should be associated with this host in the <a href="#">statusmap</a> CGI. You can specify a JPEG, PNG, and GIF image if you want, although I would strongly suggest using a GD2 format image, as other image formats will result in a lot of wasted CPU time when the statusmap image is generated. GD2 images can be created from PNG images by using the <b>pngtogd2</b> utility supplied with Thomas Boutell's <a href="#">gd library</a> . The GD2 images should be created in <i>uncompressed</i> format in order to minimize CPU load when the statusmap CGI is generating the network map image. The image will look best if it is 40x40 pixels in size. You can leave these option blank if you are not using the statusmap CGI. Images for hosts are assumed to be in the <b>logos/</b> subdirectory in your HTML images directory (i.e. <code>/usr/local/nagios/share/images/logos</code> ). |

- 2d\_coords:** This variable is used to define coordinates to use when drawing the host in the [statusmap](#) CGI. Coordinates should be given in positive integers, as the correspond to physical pixels in the generated image. The origin for drawing (0,0) is in the upper left hand corner of the image and extends in the positive x direction (to the right) along the top of the image and in the positive y direction (down) along the left hand side of the image. For reference, the size of the icons drawn is usually about 40x40 pixels (text takes a little extra space). The coordinates you specify here are for the upper left hand corner of the host icon that is drawn. Note: Don't worry about what the maximum x and y coordinates that you can use are. The CGI will automatically calculate the maximum dimensions of the image it creates based on the largest x and y coordinates you specify.
- 3d\_coords:** This variable is used to define coordinates to use when drawing the host in the [statuswrl](#) CGI. Coordinates can be positive or negative real numbers. The origin for drawing is (0.0,0.0,0.0). For reference, the size of the host cubes drawn is 0.5 units on each side (text takes a little more space). The coordinates you specify here are used as the center of the host cube.

## Extended Service Information Definition

### Description:

Extended service information entries are basically used to make the output from the [status](#) and [extinfo](#) CGIs look pretty. They have no effect on monitoring and are completely optional.

### Definition Format:

Note: Variables in red are required, while those in black are optional. However, you need to supply at least one optional variable in each definition for it to be of much use.

```
define serviceextinfo{
    host_name          host_name
    service_description service_description
    notes              note_string
    notes_url          url
    action_url         url
    icon_image         image_file
    icon_image_alt     alt_string
}
```

### Example Definition:

```
define serviceextinfo{
    host_name          linux2
    service_description Log Anomalies
    notes              Security-related log anomalies on secondary Linux server
    notes_url          http://webserver.localhost.localdomain/serviceinfo.pl?host=linux2&service=Log+Anomalies
    icon_image          security.png
    icon_image_alt     Security-Related Alerts
}
```

## Variable Descriptions:

- host\_name:** This directive is used to identify the *short name* of the host that the [service](#) is associated with.
- service\_description:** This directive is description of the [service](#) which the data is associated with.
- notes:** This directive is used to define an optional string of notes pertaining to the service. If you specify a note here, you will see the it in the [extended information](#) CGI (when you are viewing information about the specified service).
- notes\_url:** This directive is used to define an optional URL that can be used to provide more information about the service. If you specify an URL, you will see a link that says "Extra Service Notes" in the [extended information](#) CGI (when you are viewing information about the specified service). Any valid URL can be used. If you plan on using relative paths, the base path will the the same as what is used to access the CGIs (i.e. */cgi-bin/nagios/*). This can be very useful if you want to make detailed information on the service, emergency contact methods, etc. available to other support staff.
- action\_url:** This directive is used to define an optional URL that can be used to provide more actions to be performed on the service. If you specify an URL, you will see a link that says "Extra Service Actions" in the [extended information](#) CGI (when you are viewing information about the specified service). Any valid URL can be used. If you plan on using relative paths, the base path will the the same as what is used to access the CGIs (i.e. */cgi-bin/nagios/*).
- icon\_image:** This variable is used to define the name of a GIF, PNG, or JPG image that should be associated with this host. This image will be displayed in the [status](#) and [extended information](#) CGIs. The image will look best if it is 40x40 pixels in size. Images for hosts are assumed to be in the **logos/** subdirectory in your HTML images directory (i.e. */usr/local/nagios/share/images/logos*).
- icon\_image\_alt:** This variable is used to define an optional string that is used in the ALT tag of the image specified by the *<icon\_image>* argument. The ALT tag is used in the [status](#), [extended information](#) and [statusmap](#) CGIs.
-

# External Command File Permissions

---

## Notes

These instructions assume that you've installed Nagios on a dedicated monitoring/admin box that doesn't contain normal user accounts (i.e. isn't a public machine). If you've installed Nagios on a public/multi-user machine, I would suggest setting more restrictive permissions on the external command file and using something like [CGIWrap](#) to run the CGIs as a specific user. Failing to do so may allow normal users to control Nagios through the external command file! I'm guessing you don't want that. More information on securing Nagios can be found [here](#).

## Introduction

One of the most common problems people have seems to be with setting proper permissions for the external command file. You need to set the proper permission on the `/usr/local/nagios/var/rw` **directory** (or whatever the path portion of the `command_file` directive in your [main configuration file](#) is set to). I'll show you how to do this. Note: You must be *root* in order to do some of these steps...

## Users and Groups

First, find the user that your web server process is running as. On many systems this is the user *nobody*, although it will vary depending on what OS/distribution you are running. You'll also need to know what user Nagios is effectively running as - this is specified with the `nagios_user` variable in the main config file.

Next we're going to create a new group whose members include the user the web server is running as and the user Nagios is running as. Let's say we call this new group '**nagiocmd**' (you can name it differently if you wish). On RedHat Linux you can use the following command to add a new group (other systems may differ):

```
/usr/sbin/groupadd nagiocmd
```

Next, add the web server user (*nobody* or *apache*, etc) and the Nagios user (*nagios*) to the newly created group with the following commands:

```
/usr/sbin/usermod -G nagiocmd nagios  
/usr/sbin/usermod -G nagiocmd nobody
```

## Creating the directory

Next, create the directory where the command file should be stored. By default, this is `/usr/local/nagios/var/rw`, although it can be changed by modifying the path specified in the `command_file` directive.

```
mkdir /usr/local/nagios/var/rw
```

## Setting directory permissions

Next, change the ownership of the directory that will be used to hold the command file...

```
chown nagios.nagiocmd /usr/local/nagios/var/rw
```

Make sure the Nagios user has full permissions on the directory...

```
chmod u+rwX /usr/local/nagios/var/rw
```

Make sure the group we created has full permissions on the directory.

```
chmod g+rwX /usr/local/nagios/var/rw
```

In order to force newly created files in the directory to inherit the group permissions from the directory, we need to enable the group sticky bit on the directory...

```
chmod g+s /usr/local/nagios/var/rw
```

### Verifying the permissions

Check the permissions on the `rw/` subdirectory by running `'ls -al /usr/local/nagios/var'`. You should see something similar to the following:

```
drwxrws---  2 nagios nagiocmd    1024 Aug 11 16:30 rw
```

Note that the user *nagios* is the owner of the directory and the group *nagiocmd* is the group owner of the directory. The *nagios* user has **rwX** permissions and group *nagiocmd* has **rw** permissions on the directory. Also, note that the group sticky bit is enabled. That's what we want...

### Restart your web server

Once you set the proper permission on the directory containing the external command file, make sure to restart your web server. If you fail to do this, Apache will not be able to write to the external command file, even though the user it runs as is a member of the *nagiocmd* group.

### Additional notes...

If you supplied the `--with-command-grp=somegroup` option when running the configure script, you can create the directory to hold the command file and set the proper permissions automatically by running `'make install-commandmode'`.

---

# Extended Information Configuration

---

## What is Extended Information?

Extended information consists of *optional* definitions for hosts and services that is used by the CGIs in the following ways:

- to provide URLs to additional information about the host or service
- to add pretty icons to the hosts and services displayed in the web interface
- to draw hosts in the [statusmap](#) and [statuswrl](#) CGIs at user-defined 2-D and 3-D coordinates

## Where is Extended Information Defined?

Extended information definitions are stored in [object configuration files](#) along with definitions for hosts, services, contacts, etc. You can use templates to define entries for multiple hosts and services quickly and easily.

---